



VNIVERSITAT
DE VALÈNCIA

School of Psychology

How Literacy Shapes Orthographic Processing

A Doctoral Dissertation with International Doctorate
Mention

Presented in Partial Fulfillment of the Requirements
for the Degree of Doctor of Philosophy:

Reading and Comprehension

by

María Fernández-López

València, February 2022

Supervisors

Dr. Manuel Perea Lara*

Dr. Marta Vergara-Martínez

Dr. Ana Marcet Herranz

*Thesis tutor



València, 6 de febrero de 2022

Los doctores Manuel Perea Lara (Universitat de València), Marta Vergara Martínez (Universitat de València) y Ana Marcet (Universitat de València)

DECLARAN:

Que el trabajo titulado *How literacy shapes orthographic processing*, que presenta María Fernández López para la obtención del título de doctora, se ha realizado bajo nuestra dirección y tutoría cumple todos los requisitos para poder optar a su lectura como Tesis Doctoral con Mención Internacional en la Universitat de València.

Y para que así conste y tenga los efectos oportunos, firmamos el presente documento.

Manuel Perea Lara

Marta Vergara Martínez

Ana Marcet Herranz

Para la realización de la Tesis Doctoral Internacional presentada, la investigadora María Fernández-López ha disfrutado de:

- La ayuda para contratos predoctorales para la formación de doctores del Ministerio de Ciencia, Innovación y Universidades (PRE2018-083922)

La presente Tesis Doctoral Internacional se engloba dentro de los siguientes proyectos de investigación:

- (2021-2024) Aprende, entrena, lee: Una aproximación multidimensional al procesamiento ortográfico (Ministerio de Ciencia, e Innovación; código PID2020-116740GB-I00). Universitat de València.
- (2018-2021) Descifrando los mecanismos del léxico mental: Desde el aprendizaje ortográfico hasta la integración semántica (Ministerio de Ciencia, Innovación y Universidades.; código PSI2017-86210-P). Universitat de València.

En este periodo y para difundir el trabajo de investigación recogido en esta Tesis Doctoral, la investigadora María Fernández-López ha obtenido:

- Una beca de la Psychonomic Society (Graduate Student Travel Award) en el año 2021
- Una beca de la Cognitive Neuroscience Society (Graduate Student Award) en el año 2021
- Una Beca de la Sociedad Experimental de Psicología Experimental (SEPEX) para la Difusión de Trabajos de Investigación. Convocatoria 2018-2019

Para la obtención de la Tesis Doctoral presentada con mención de “Tesis Internacional”, la investigadora ha realizado dos estancias internacionales en:

- (2019) Laboratoire de Psychologie Cognitive, Aix-Marseille Université, Marseille (Francia). Estancia financiada por el Programa Erasmus+ de la UE.

- (2021) Language, Learning and Reading Lab, Scuola Internazionale Superiore di Studi Avanzati, Trieste (Italia)

Status of the articles

Fernández-López, M., Marcet, A., & Perea, M. (2019). Can response congruency effects be obtained in masked priming lexical decision? *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 45, 1683–1702. doi:10.1037/xlm0000666

(Scopus: Q1 [Experimental and Cognitive Psychology]; JCR: Factor de impacto 3.051)

Fernández-López, M., Gómez, P., & Perea, M. (2021). Which factors modulate letter position coding in pre-literate children? *Frontiers in Psychology*, 12, 708274. doi:10.3389/fpsyg.2021.708274

(Scopus: Q2 [Psychology (Miscellaneous)]; JCR: Factor de impacto 2.990)

Fernández-López, M., Marcet, A., & Perea, M. (2021). Does orthographic processing emerge rapidly after learning a new script? *British Journal of Psychology*, 112, 52–91. doi:10.1111/BJOP.12469

(Scopus: Q1 [Psychology (Miscellaneous)]; JCR: Factor de impacto 4.267)

Fernández-López, M., Mirault, J., Grainger, J., & Perea, M. (2021). How resilient is reading to letter rotations? *A parafoveal preview investigation. Journal of Experimental Psychology: Learning, Memory, and Cognition*, 47, 2029–2042. doi:10.1037/xlm0000999

(Scopus: Q1 [Experimental and Cognitive Psychology]; JCR: Factor de impacto 3.051)

Fernández-López, M., Davis, C. J., Perea, M., Marcet, A., & Gómez, P. (in press). Unveiling the boost in the sandwich priming technique. *Quarterly Journal of Experimental Psychology*. doi:10.1177/17470218211055097

(Scopus: Q1 [Experimental and Cognitive Psychology]; JCR: Factor de impacto 2.143)

Manuscripts sent for peer review at the time of the present thesis submission

Fernández-López, M. & Perea, M. A letter is a letter and its co-occurrences:
cracking the emergence of location-invariant processing

Fernández-López, M., Perea, M. & Vergara-Martínez, M. On the time course of
the resilience of letter detectors to rotations: A masked priming ERP
investigation

Fernández-López, M., Gomez, P. & Perea, M. Letter rotations: Through the
Magnifying Glass and What Evidence Found There

A miña familia

A la meua família

Agradecimientos

Pero, ao cabo, un non leva nada máis fondo e permanente que as súas vivencias e recordos abandonados: a néboa de cinco días, a friaxe da nosa infancia, o aprendido durante a formación da alma. E, máis tarde, un vaise alimentando ademáis do lido, as súas viaxes e experiencias pessoais, que poden enriquecernos ou perdernos nun real extravío sen raíces, pero que soen filtrar as augas estancadas dunha pobre indiferencia ao que nos é máis próximo.

Lois Pereiro

Nesta tese falo da alfabetización como unha ferramenta que modifica a maquinaria cerebral para poder sostener un proceso tan complicado como é a lectura. De maneira similar, as persoas das que eliximos arrodearnos son as pezas das engrenaxes que fan xirar un proceso tan complicado como é a vida.

A mamá, a papá e a Manu. O que máis quero no mundo. A miña casiña e o meu lar. Ensinástesme o esforzo e sodes agora o acougo. Síntovos sempre moi preto.

Muchas gracias a todas las personas que me han acompañado, de un modo y otro, en la Universitat de València. Gracias a todas las personas que han supervisado mi trabajo, que no han sido pocas. No estaría escribiendo esto si no fuese por Manolo. Gracias por confiar en mí. Por emplear tu tiempo en enseñarme una parte de todo lo que sabes. Por tu eterna disponibilidad, acompañamiento y consejo. Por apretar e impulsar. Y gracias, como no, por esas comidas. Para mí eres un MENTOR, en mayúsculas. Gracias a Marta, por tu comprensión y tus ánimos. Por enseñarme la ciencia desde otra perspectiva. Y por demostrarme que las cosas de palacio van despacio, pero al final, van. Muchas gracias a Pablo y a Joana, que habéis recibido más de un correo mío con el borrador de un manuscrito y me lo habéis devuelto con comentarios y sugerencias. De vosotros también he aprendido mucho. Y gracias a Baciero, te elegiría para subir a cualquier cima.

A la meua família del País Valencià. A totes les persones amb les que he compartit el pa, l'aigua i les cases. M'heu trencat el cor, dividit ara entre dues llengües i dues terres. Soc qui soc gràcies també a la vostra amistat. Xiques, nínas, no tinc paraules per descriure el que sent per vosaltres. Encara sort que no em calen. Sou les cures, la reflexió i l'amor portat a l'infinit. Vos estime molt i vos estic molt (molt!) agraïda, per tot. Vicentet, amic, moltes gràcies també a tu, per la complicitat i l'humor.

Gracias a Sara, por tu cálido (calidísimo) trabajo. No habría llegado tan entera si no fuera por ti, ya sea por el colchón que has supuesto estos años o por la fuerza que me has ayudado a construir.

Grazas a Andrés e ao Vila, por transmitirme sempre o voso orgullo e a vosa enerxía. Ocupades un cachiño no meu corazón. E sodes ámbolos dous un exemplo a seguir para min, cada un ao seu xeito, como profesionais e como persoas. Quérovos ao meu carón.

Hablando de orgullo, me vienes a la cabeza, Luchiño, mano a mano nos sacamos una carrera. Y me acuerdo de ti siempre que me pongo delante de una clase. Moitas grazas tamén a Marti, pola túa escoita e apoio incondicional, e ao resto das pepis, por celebrarmos xuntas os pasos adiante.

Special affection and gratitude to all the people that Academia has given me. People who have encouraged critical thinking, who have made science fun and who, without realizing it, have reinvigorated me. A *merci spéciale* to the laboratory of Jonathan Grainger and a *grazie speciale* to the Crepaldi's lab.

Finalmente, muchas gracias a todas las personas que se esfuerzan día a día para que el personal investigador tenga derechos laborales.

How Literacy Shapes Orthographic Processing

Content

Resumen	7
La aparición del procesamiento ortográfico	8
La resiliencia de los detectores de letras	13
La contribución metodológica a la investigación del reconocimiento de palabras	18
Introduction	27
The emergence of orthographic processing	31
Which factors modulate letter position coding in pre-literate children?	35
Abstract	35
Introduction	36
Method	41
Results	42
Discussion	44
Does orthographic processing emerge rapidly after learning a new script?	49
Abstract	49
Introduction	50
Experiment 1: location-invariant processing	59
Method	59
Results	65
Discussion	71
Experiment 2: location-specific processing	72
Method	72
Results	74
Discussion	77
General Discussion	77
A letter is a letter and its co-occurrences: Cracking the emergence of location-invariance processing	101
Abstract	101
Introduction	102
Method	108
Results	111
Discussion	113
Conclusions	121

<i>The resilience of letter detectors.....</i>	125
How resilient is reading to letter rotations? A parafoveal preview investigation.....	129
Abstract	129
Introduction.....	130
Method.....	135
Results	138
Discussion.....	144
On the time course of the resilience of letter detectors to rotations: A masked priming ERP investigation.....	149
Abstract	149
Introduction.....	150
Method.....	158
Results	162
Discussion.....	168
Letter rotations: Through the Magnifying Glass and What Evidence Found There.....	171
Abstract	171
Introduction.....	172
Experiment 1: Lexical Decision Task.....	178
Method	178
Results and Discussion.....	180
Experiment 2: Semantic Categorization Task.....	185
Method	185
Results and Discussion.....	186
Experiment 3: Semantic Categorization Task.....	189
Method	189
Results and Discussion.....	190
General Discussion	192
Conclusions	195
<i>Methodological contributions to the visual word recognition research.....</i>	199
Can response congruency effects be obtained in masked priming lexical decision?	203
Abstract	203
Introduction.....	204
Experiment 1.....	213
Method	213
Results and discussion	214

Experiment 2.....	218
Method	218
Results and discussion	219
Experiment 3.....	229
Method	229
Results and discussion	229
General Discussion	233
Unveiling the boost in the sandwich priming technique	237
Abstract	237
Introduction.....	238
Method.....	244
Results	246
Discussion	253
Conclusions	259
<i>Concluding Remarks.....</i>	263
<i>References.....</i>	269
Resumen.....	269
Introduction.....	274
The emergence of the orthographic processing	274
The resilience of letter detectors.....	287
Methodological contribution to the visual word recognition research.....	297
<i>Appendix: journal version of papers</i>	309

Resumen

La lectura es una habilidad indispensable en nuestra sociedad. Sin ella, no podríamos acceder a la historia, o aprender de nuestros avances y retrocesos. Resulta casi milagroso que el Homo sapiens hiciera unas marcas en una piedra para aumentar drásticamente su capacidad de memoria. Estas marcas en la piedra, que serían lo que ahora llamamos letras, se convirtieron en herramientas. Herramientas que alteran la maquinaria cerebral para conectar las formas visuales con un significado, es decir, para sustentar la lectura.

Hoy en día, la lectura es parte de nuestra vida cotidiana hasta tal punto que no nos damos cuenta de su complejidad. Cuando leemos, nuestros ojos saltan de palabra a palabra, fijándose en cada una de ellas (normalmente) solo una vez. Es en este momento cuando sucede uno de los pasos esenciales de la lectura eficiente: el reconocimiento visual de la palabra. Cuando leemos una palabra, se producen una serie de complejos procesos cognitivos que implican operaciones en diferentes niveles de procesamiento (perceptivo, ortográfico, fonológico, morfológico, sintáctico y semántico) durante unos cientos de milisegundos. De manera importante, si hay un aspecto del reconocimiento visual de palabras sobre el que existe un consenso general entre las personas que investigan la lectura, ese es el papel crucial de las letras. Por lo general, se asume que el reconocimiento de una palabra escrita está mediado por el procesamiento de las letras que la componen, al menos en las lenguas que utilizan la ortografía alfabética. La presente tesis aspira a ser una contribución más al corpus de investigación de estas "células" de la lectura. En concreto, el objetivo general es examinar los procesos que se sitúan en el puente entre las etapas de procesamiento visual, de bajo nivel, y el procesamiento, de alto nivel, de las palabras: el procesamiento ortográfico.

El procesamiento ortográfico es la codificación de la información sobre las identidades y la posición de las letras; aquello que nos permite diferenciar PERRO de CERRO o CONEJO de ENCOJO. Para avanzar en nuestra comprensión de la

lectura, esta tesis se centra en cómo la alfabetización moldea la cognición para dar lugar a estos procesos cruciales del reconocimiento visual de palabras. Para ello, la tesis se estructura en tres secciones: (i) la emergencia del procesamiento ortográfico, (ii) la resiliencia de los detectores de letras, y (iii) las contribuciones metodológicas a la investigación del reconocimiento visual de palabras. El bloque sobre la emergencia del procesamiento ortográfico comprende tres estudios que examinan el desarrollo del procesamiento ortográfico con la alfabetización. En el bloque sobre la resiliencia de los detectores de letras, presento tres estudios que examinaron cómo la distorsión visual afecta al procesamiento ortográfico. Para ello, nos centramos en la distorsión visual en una única dimensión: la rotación de letras. Finalmente, la última sección incluye una contribución metodológica a la investigación sobre el reconocimiento de palabras. En concreto, comprende dos estudios realizados para entender mejor los fundamentos cognitivos del *masked priming*, la herramienta metodológica de referencia para estudiar los primeros momentos del reconocimiento visual de palabras.

La aparición del procesamiento ortográfico

Aprender a leer implica un ajuste de nuestro cerebro de primate. Cuando aprendemos a leer, nuestro cerebro recicla algunas regiones del córtex visual para asumir una nueva tarea, la del reconocimiento de palabras. Uno de los hitos centrales de este proceso es la aparición del procesamiento ortográfico. El aprendizaje de la codificación de las identidades de las letras, entendida como el mapeo de la entrada visual a las representaciones abstractas de las letras (*A* a **a** a comparten la misma representación abstracta, independientemente de su forma escrita) ocurriría en los dos primeros años de la adquisición de la lectura (Jackson y Coltheart, 2001; Gómez y Perea, 2020). Sin embargo, aún no se entiende claramente cuándo y cómo surge la codificación de la posición de las letras.

Determinar los factores que influyen en la aparición y el desarrollo del procesamiento ortográfico es fundamental, ya que nos permite identificar a los

niños y niñas que necesitan una intervención previa a la alfabetización y, en consecuencia, minimizar posibles dificultades de lectura. Con este propósito, realizamos el estudio titulado *¿Qué factores modulan la codificación de la posición de las letras en los niños prealfabetizados?* (Which Factors Modulate Letter Position Coding in Pre-literate Children? Fernández-López et al., 2021, *Frontiers in Psychology*). Nuestro objetivo fue indagar en los precursores del mecanismo responsable de la codificación del orden de las letras, aquel que nos permite diferenciar CONEJO de ENCOGER. Sabemos que los mejores lectores codifican el orden de las letras con mayor precisión que los peores lectores (véase Gómez et al., 2021; Pagán et al., 2021), por lo tanto, es crucial arrojar luz sobre las raíces de ese procesamiento. Para ello, examinamos los posibles vínculos entre la codificación del orden de las letras en infantes pre-alfabetizados con los procesos cognitivos básicos relacionados con la adquisición temprana de la lectura (atención, sensación, percepción y memoria), evaluados con la Batería de Inicio a la Lectura (Sellés et al., 2016). Para medir la codificación del orden, se llevó a cabo una tarea de emparejamiento igual-diferente, en la que los niños y niñas tuvieron que decidir si dos pares de letras eran iguales (TZ-TZ) o diferentes (con letras traspuestas: TZ-ZT) (véase Perea et al., 2016). Los resultados de los análisis correlacionales mostraron que los niños y niñas pre-alfabetizados que mejor diferenciaban los pares de letras traspuestas (TZ-ZT) de los pares de letras idénticas (TZ-TZ) eran los que tenían puntuaciones más altas en las pruebas de procesos cognitivos básicos. Así pues, las habilidades cognitivas básicas se asocian con la capacidad de codificar el orden serial en las cadenas de letras.

La investigación con personas lectoras expertas ha demostrado que, una vez que se domina la lectura, la categoría de las letras evoluciona de ser un mero objeto visual a gozar de un estatus específico. Por ejemplo, cuando se presenta una cadena de símbolos (£\$?%@) las personas fijan su atención en el centro de la cadena. Sin embargo, cuando se presenta una cadena de letras (FGJGM), se fijan también en las letras exteriores, especialmente en la primera letra (Tydgate y Grainger, 2009). Este fenómeno se debe al procesamiento específico de la localización, responsable de la codificación de la posición de las letras. Este

procesamiento nace de una adaptación de los mecanismos de la codificación visual de objetos a las particularidades del procesamiento de palabras.

Además, la codificación del orden de las letras se vuelve más flexible, lo que conduce al procesamiento invariante de la localización; *FGJM* se confunde fácilmente con el *FJGM*, del mismo modo, *cuenro* se confunde fácilmente con *cuerno* y el *cholocate* con el *chocolate*. Este fenómeno se suele captar cuando se utiliza una tarea de emparejamiento igual-diferente que consiste en decidir si dos pares de estímulos son iguales (*FGJM-FGJM*) o no (*FGJM-FPCM*). La evidencia muestra que las respuestas "no" a un par con letras transpuestas (*FGJM-FJGM*) son más lentas y más propensas a errores que las respuestas "no" a un par con letras remplazadas (*FGJM-FPCM*). Este efecto, conocido como efecto de transposición (TL), también se produce con cadenas compuestas por símbolos, números o letras desconocidas (Duñabeitia et al., 2012; Massol et al., 2013). No obstante, lo curioso es que es sustancialmente mayor para las cadenas de letras conocidas. La disociación entre letras y otros objetos visuales se explica en la existencia de un mecanismo ortográfico específico, que se utiliza para codificar el orden invariante de las letras en la palabra, y opera complementando la incertidumbre perceptiva sobre qué posición ocupa un objeto visual en una serie (Grainger, 2018; véase también Marcet et al., 2019; Massol et al., 2013). De manera importante, se ha postulado que este mecanismo ortográfico emerge rápidamente en los primeros estadios de adquisición de la alfabetización (Dandurand et al., 2010; Duñabeitia et al., 2014, 2015).

Sin embargo, los datos empíricos sobre la aparición del procesamiento invariante de localización son muy escasos y no concluyentes. La razón es que no es sencillo comparar el rendimiento de los niños y niñas antes y después de adquirir la alfabetización, dado que hay mucha variabilidad entre grupos, lo que conduce a problemas de interpretación (véase Perea et al., 2016). Una opción para superar esta limitación es diseñar un análogo del aprendizaje de la lectura con personas adultas aprendiendo un nuevo alfabeto. Esta fue la estrategia seguida en el segundo y tercer estudio. En el segundo estudio, *¿Surge rápidamente el procesamiento ortográfico tras el aprendizaje de un nuevo*

alfabeto? (Does Orthographic Processing Emerge Rapidly After Learning a New Script?, Fernández-López et al., 2021, *British Journal of Psychology*), comprobamos si el procesamiento ortográfico, examinado a través del procesamiento invariante de localización y el procesamiento específico de localización, emergía rápidamente con la alfabetización. Para ello, diseñamos un estudio en el que se enseñó a personas adultas a leer y escribir en un alfabeto artificial durante seis sesiones. Además, se familiarizaron visualmente con un alfabeto artificial diferente, que se utilizó como control. En lo que respecta al procesamiento de localización invariante, realizamos tareas de emparejamiento igual-diferente en las que los ensayos diferentes estaban compuestos por pares con letras traspuestas (FGJM-FJGM) y pares con letras remplazadas (FGJM-FPCM). Las tareas se realizaron con los alfabetos entrenado y control, antes y después de aprender a leer. Los resultados mostraron efectos de trasposición similares en los alfabetos entrenado y control en ambas fases, pre- y post-aprendizaje. De ello concluimos que el aprendizaje de la lectura y la escritura en un nuevo alfabeto no origina la rápida aparición del procesamiento invariante de la localización. Sin embargo, las respuestas a los pares iguales (FGJM-FGJM) fueron más rápidas y más precisas en la fase de post-entrenamiento para el alfabeto entrenado, pero no para el alfabeto control. Por lo tanto, parece que la alfabetización sí indujo una especialización rudimentaria para las letras, que facilitó la identificación de los pares de letras iguales. Para comprobar la aparición del procesamiento específico de la localización, llevamos a cabo tareas de identificación de caracteres con los alfabetos entrenado y control antes y después de aprender a leer. En este tipo de tareas, se presenta muy brevemente una cadena de cinco caracteres (FGJGM) mientras la persona mira el centro de la cadena. La cadena se sigue de una máscara con una señal que indica una de las posiciones y dos alternativas de letras (####; alternativas G y L). La tarea consiste en elegir, de entre las dos alternativas, la que coincida con la identidad la letra en el lugar indicado (G→FGJGM). Los resultados mostraron una clara ventaja de la posición media para ambos alfabetos, patrón característico cuando se emplean símbolos, pero no letras. Por lo tanto, aprender a leer y escribir en

un nuevo alfabeto tampoco facilitó a la rápida aparición del procesamiento específico de la localización.

El procesamiento ortográfico necesita, probablemente, mucho más tiempo de práctica lectora para surgir. Una alternativa es que podría requerir de experiencia con estructuras ortográficas específicas. En el tercer estudio de la presente sección, *Una letra es una letra y sus co-ocurrencias: Descifrando la emergencia del procesamiento invariante de localización* (A letter is a letter and its co-occurrences: cracking the emergence of location-invariant processing; enviado para revisión por pares), pusimos a prueba si el procesamiento invariante de localización surgía rápidamente tras la exposición a las regularidades ortográficas—bigramas¹—en un alfabeto nuevo. Para ello, diseñamos un estudio con dos fases. En la fase 1 replicamos el protocolo seguido por Chetail (2017). Concretamente, las personas fueron expuestas a un flujo de palabras con letras artificiales durante unos minutos, con cuatro bigramas apareciendo con frecuencia (cada palabra contenía uno de esos cuatro bigramas en una posición específica; **QD**⌘△, **L**Γ⋄△◁, **6**ΠΤ⊥△, ◁Π⌘⊕△, la negrita indica el bigrama frecuente). A continuación, se llevó a cabo una tarea en la que las personas tenían que decidir, entre dos palabras con letras artificiales, cuál habían visto con anterioridad. Los resultados mostraron que las palabras con bigramas entrenados fueron juzgadas como más parecidas a las vistas, por lo que las personas adquirieron rápidamente las nuevas regularidades ortográficas. En la fase 2, se realizó una tarea de emparejamiento igual-diferente para examinar las posibles diferencias entre pares con letras traspuestas en un bigrama frecuente (entrenado) y pares con letras traspuestas en un bigrama no frecuente (no entrenado). Los resultados mostraron que las respuestas a los pares con letras traspuestas en un bigrama frecuente eran menos precisas que las respuestas a pares con las letras traspuestas en un bigrama infrecuente. Este hallazgo sugiere que la exposición repetida a regularidades ortográficas

¹ Un bigrama es la co-ocurrencia de dos letras en una palabra. La palabra árbol se compone de diez bigramas: AR-AB-AO-AL-RB-RO-RL-BO-BL-OL

facilitaría el desarrollo de representaciones internas de grupos de letras (bigramas), induciendo el procesamiento invariable de localización.

En resumen, el objetivo de los tres estudios de esta sección fue ampliar nuestra comprensión acerca de los fundamentos del procesamiento ortográfico. Cabe destacar las siguientes premisas: (i) La preparación cognitiva para la lectura surge en una fase temprana del desarrollo, con la adaptación de los mecanismos de procesamiento de los objetos visuales. Por ello, (ii) la detección temprana de las deficiencias en el análisis visual del *input* es crucial para prevenir posibles dificultades lectoras. Además, hemos añadido a la literatura que (iii) el procesamiento ortográfico no emerge rápidamente cuando se aprende a leer. No obstante, como preludio, (iv) algunos procesos no especializados se ponen a punto para guiar el progreso funcional de la lectura.

Por lo tanto, teniendo en cuenta todo esto, la enseñanza del aprendizaje de la lectura no debería basarse únicamente en el conocimiento de las letras y la decodificación fonológica. Los y las profesionales la enseñanza deben considerar que la lectura se construye sobre un fundamento cognitivo básico que puede entrenarse para facilitar el aprendizaje lector. Por citar algunas sugerencias: el entrenamiento en una tarea de emparejamiento igual-diferente con estímulos de letras puede ayudar a la codificación del orden de éstas; además, la mera exposición visual a las palabras ayuda a crear una capa rudimentaria de bigramas que facilitarían la posterior codificación de las palabras.

La resiliencia de los detectores de letras

Con la adquisición de la lectura siguiendo su curso, se construye un sistema de decodificación flexible y potente: las personas lectoras pueden acceder rápidamente a las representaciones abstractas de las letras a pesar de su formato visual (*a A a a* comparten la misma representación; véase Grainger y Dufau, 2012). Sin embargo, la evidencia empírica sobre la resiliencia del sistema de reconocimiento de palabras a la distorsión perceptiva es todavía escasa,

probablemente porque conlleva una deformación en dimensiones diferentes. De esta manera, cualquier diferencia de procesamiento entre una palabra deformada y una en el formato prístino que solemos encontrar al leer podría deberse a múltiples factores (véase Vergara-Martínez et al., 2021). En esta sección, presento tres estudios que examinaron la tolerancia del sistema de reconocimiento de palabras a la distorsión visual en una sola dimensión: las rotaciones de las letras.

Un valor añadido de estudiar la rotación es que dos de los principales modelos de reconocimiento visual de palabras hacen predicciones explícitas sobre cómo las rotaciones de las letras afectan a la lectura de palabras. El modelo SERIOL de reconocimiento de palabras (Whitney, 2001, 2002) asume que "los niveles de entrada a las unidades de letras se reducen para los estímulos rotados" (p. 119, Whitney, 2002), aumentando gradualmente el tiempo necesario para codificar una cadena de letras. En cambio, el modelo de los detectores de combinación local (LCD; Dehaene et al., 2005) ofrece una explicación de todo-o-nada sobre cómo la rotación de las letras afecta al procesamiento de las mismas. Postula un límite de 40°-45°, más allá del cual los detectores de letras no serían capaces de hacer frente a las rotaciones de manera eficiente. No obstante, habría poco coste para las rotaciones de letras por debajo de ese límite.

La premisa de un coste de rotación de las letras ha sido puesta en duda recientemente por una serie de estudios que han demostrado que las lectoras pueden hacer frente a rotaciones de hasta 180° desde el acceso inicial a las representaciones ortográficas. Por ejemplo, con la técnica de *priming* enmascarado, Perea et al. (2018) encontraron efectos de identidad considerables para palabras rotadas 90°, comparables a los obtenidos con estímulos canónicos (0°). Este patrón de resultados se extendió a 180° en un estudio realizado por Yang y Lupker (2019). Sin embargo, se puede argumentar que, como ambos, *prime* y estímulo-test, estaban rotados, las personas participantes podrían haber desarrollado estrategias conscientes, como inclinar la cabeza, para hacer frente a las rotaciones. Por lo tanto, una estrategia adecuada sería presentar los estímulos objetivo en la orientación canónica,

precedidas por los *primes* rotados. Esta fue la estrategia seguida por Benyhe y Csibri (2021), que encontraron un considerable efecto de *priming* enmascarado para ángulos de hasta 90°.

Cabe destacar en este punto que, cuando una palabra está completamente rotada, las personas pueden rotar mentalmente el estímulo como un todo, con un solo movimiento de rotación mental. En cambio, si rotamos las letras individuales dentro de las palabras, los individuos necesitarían más de un movimiento de rotación. Así, esta manipulación nos permite examinar estrictamente las predicciones de los modelos. Por esta razón, en los tres estudios de esta sección, rotamos las palabras letra por letra. Nuestro objetivo principal era examinar la resistencia (y resiliencia) de los detectores de letras a diferentes ángulos de rotación y, de esta manera, probar las predicciones del modelo SERIOL (Whitney, 2001, 2002) y del modelo LCD (Dehaene et al., 2005).

En el primer estudio, *¿Cuán resiliente es la lectura a las rotaciones de las letras? Una investigación de previsualización parafoveal* (How Resilient is Reading to Letter Rotations? A Parafoveal Preview Investigation, Fernández-López et al., 2021, *Journal of Experimental Psychology: Learning, Memory, and Cognition*), pusimos a prueba la resiliencia de los detectores de letras a las rotaciones durante la lectura normal para examinar las predicciones de los modelos LCD y SERIOL. Llevamos a cabo un experimento de lectura con frases escritas en letras mayúsculas utilizando el paradigma de cambio de límites contingente a la mirada de Rayner (1975), en el que examinamos si las previsualizaciones parafoveales de identidad eran más efectivas que las previsualizaciones parafoveales no-relacionadas. Las letras de las previsualizaciones estaban rotadas 15°, 30°, 45° o 60°. Es importante destacar que solo las previsualizaciones parafoveales estaban rotadas. El resto del texto, incluido las palabras-test, se presentó con las letras en su orientación canónica vertical. Encontramos un patrón de resultados notablemente similar para todas las medidas de fijación que evaluamos: la duración de la primera fijación, la duración de la fijación única y la duración de la mirada. La ventaja de la condición de identidad fue mayor cuando las letras de la previsualización se rotaron 15°, y fue numéricamente menor, pero aún robusta, a 30°. Para las rotaciones de 45°,

la ventaja de la condición de identidad fue sólo marginal. Por último, para las rotaciones de 60°, no hubo indicios de beneficio de la condición de identidad. Estos resultados sugieren que, aunque los detectores de letras parecen estar gravemente afectados por ángulos de rotación de 45° o más en los primeros momentos del procesamiento, como postula el modelo LCD, también hubo un coste gradual para las rotaciones de letras por debajo de ese límite, como postula el modelo SERIOL.

En el segundo estudio, *Sobre el curso temporal de la resiliencia de los detectores de letras a las rotaciones: Una investigación con priming enmascarado y potenciales relacionados con eventos* (On the time course of the resilience of letter detectors to rotations: A masked priming ERP investigation; enviado para revisión por pares) examinamos el curso temporal del efecto de la rotación de las letras en un paradigma de *priming* enmascarado de identidad. Nos centramos en probar la predicción del modelo LCD: "los detectores de letras se ven dañados por la rotación ($> 40^\circ$)" (Dehaene et al., 2005, p. 340). Mientras el estímulo-test se presentó en la orientación canónica, las letras de los *primes* se presentaron en ángulos de 0°, 45° ó 90°. Los datos conductuales mostraron que el *priming* de identidad estaba modulado por la rotación de las letras: era más fuerte para los *primes* de 0°, disminuía para los de 45° y era insignificante para los de 90°. Los datos de los potenciales relacionados con eventos (PREs) nos permitieron examinar en detalle el curso temporal del procesamiento de las palabras con letras rotadas. Las comparaciones de amplitudes mostraron que el efecto de la relación entre estímulos *prime* y test surgió a los 150 ms para la orientación estándar del *prime* (0°). De hecho, el efecto de identidad para los *primes* de 0° siguió el curso natural (es decir, comenzó alrededor de 100-150 ms y se fortaleció en las siguientes etapas de procesamiento). Por el contrario, la aparición del efecto para los *primes* de 45° se retrasó, y su tamaño también se redujo. Surgió alrededor de 170-300 ms y fue más débil que para los estímulos de 0°. En el caso de los *primes* de 90°, no hubo efecto de identidad en ningún componente. Este patrón sugiere que las rotaciones de las letras en torno a 45° impiden la integración automática de la representación ortográfica de los *primes*

con los estímulos objetivo, proporcionando así pruebas que apoyan las predicciones del modelo LCD de reconocimiento de palabras.

En el tercer estudio, *Rotaciones de letras: A través de la lupa y lo que la evidencia encontró allí* (Letter rotations: Through the Magnifying Glass and What Evidence Found There; enviado para revisión por pares), llevamos a cabo tres experimentos que examinaron el coste de la rotación de las letras en el acceso al léxico. De esta manera, pusimos a prueba las predicciones de los modelos LCD y SERIOL sobre la resiliencia de los detectores de letras a ángulos de rotación relativamente moderados. En el Experimento 1, realizamos una tarea de decisión léxica con estímulos de letras rotadas 0°, 22.5°, 45° y 67.5°. En los Experimentos 2 y 3, realizamos una tarea de categorización semántica con letras rotadas 0°, 22.5°, 45° y 67.5° (Experimento 2) y 22.5°, 30°, 37.5° y 45° (Experimento 3). En todos los experimentos, las palabras podían ser de alta o baja frecuencia. Los resultados de los Experimentos 1 y 2 mostraron que, en relación con el formato canónico (0°), el coste era mínimo para las palabras de 22.5°, sustancial para las de 45° y drásticamente grande para las de 67.5°. Es importante destacar que en la tarea de decisión léxica (Experimento 1), encontramos que el efecto de la rotación de las letras era mayor para las pseudopalabras que para las palabras. Además, en el Experimento 2, descubrimos que la rotación de las letras interactuaba con la frecuencia de las palabras, pero sólo cuando las letras estaban rotadas en 67.5° (en el Experimento 1 se produjo una tendencia similar, aunque numéricamente menor). Estos hallazgos sugieren que la retroalimentación de niveles léxico-semánticos puede hacer disminuir el coste debido a la rotación de letras (véase Vergara-Martínez et al., 2021, para un argumento similar en la lectura de palabras manuscritas). Por último, en el Experimento 3 se encontró un coste gradual de lectura que aumentaba en el ángulo de 45°. En conjunto, estos resultados revelan que el coste provocado por el ángulo de rotación aumenta gradualmente en ángulos inferiores a 45°, como propone el modelo SERIOL. Los detectores de letras alcanzan un límite de resiliencia en torno a los 45° (consistente con el modelo LCD), y el coste es sustancialmente mayor en ángulos más pronunciados (67.5°).

Los tres estudios de esta sección examinaron la tolerancia del sistema de reconocimiento de palabras a la distorsión visual. Para ello, se emplearon diferentes técnicas que midieron el procesamiento de las letras rotadas: técnicas conductuales, electrofisiológicas y de registro de movimientos oculares. Teniendo en cuenta los resultados de los tres estudios y construyendo una imagen global, las principales ideas que se pueden extraer se resumen de la siguiente manera: (i) el coste de la rotación de las letras se origina muy pronto, en las primeras etapas del reconocimiento de palabras. (ii) Incluso las rotaciones moderadas suponen un coste. Sin embargo, (iii) este coste se supera fácilmente, y el sistema puede manejar rotaciones de hasta 45° en los momentos iniciales del procesamiento. Además, (iv) para hacer frente a las rotaciones mayores, los individuos parecen servirse de la retroalimentación de arriba-a-abajo (facilitación léxica). En resumen (v), encontramos que el coste provocado por el ángulo de rotación aumenta suavemente en ángulos inferiores a 45°, como propone el modelo SERIOL. Los detectores de letras alcanzan un límite de resiliencia en torno a los 45° (consistente con el modelo LCD), y el coste es muy elevado en ángulos más pronunciados. Por lo tanto, (vi) la combinación de los supuestos de los modelos SERIOL y LCD otorgaría una explicación más completa de cómo el ángulo de rotación, y más generalmente, la distorsión visual, dificulta el reconocimiento de palabras.

La contribución metodológica a la investigación del reconocimiento de palabras

Más allá de los estudios pioneros del tema (Huey, 1908), los avances más significativos en el estudio del proceso de lectura comenzaron hacia finales de la década de 1950, coincidiendo con lo que se suele llamar la revolución cognitiva (véase Pollatsek y Treiman, 2015, para una revisión). Esta revolución rompió con el enfoque conductista y dio lugar a muchas técnicas, nuevas e ingeniosas, para estudiar los procesos de lectura e identificación de palabras, incluyendo la tarea de decisión léxica (Rubinstein et al., 1970, 1971a, 1971b) y el procedimiento de *priming* enmascarado (Forster & Davis, 1984). En la tarea

de decisión léxica, se pide a los participantes que distingan lo más rápido posible las palabras de las no palabras. El tiempo que se tarda en determinar que una cadena de letras como *perro* es una palabra real proporciona un índice de lo que se tarda en contactar con la representación neural de esa palabra en nuestro diccionario interno. Si queremos enfocar las primeras etapas del procesamiento de palabras, debemos utilizar la técnica de *priming* enmascarado, considerada "la mayor innovación metodológica de las dos últimas décadas" (Grainger, 2008). En esta técnica, se presenta una máscara durante 500 ms, y luego se sustituye por un *prime* durante 50 ms, de modo que no se puede ver, que se sustituye por el estímulo-test (#####-árbol-ÁRBOL). Las personas investigadoras suelen manipular la relación prime-test, que puede estar relacionada o no (por ejemplo: #####-árbol-ÁRBOL vs. #####-panal-ÁRBOL). Las respuestas de las personas participantes suelen evaluarse con medidas conductuales, como el tiempo de reacción y la precisión, pero también con medidas más directas de la actividad cerebral, como las obtenidas con imágenes de resonancia magnética funcional (IRMf) y electroencefalograma (EEG).

En particular, la combinación de la decisión léxica y el *priming* enmascarado ha dado lugar a muchos trabajos empíricos que investigan el procesamiento de letras y palabras. Sin embargo, la técnica requiere un examen más profundo para lograr conclusiones firmes sobre los resultados obtenidos. En esta sección, se presentan dos estudios que intentan comprender los fundamentos cognitivos del *priming* enmascarado y de una de sus variantes, el *priming* en sándwich. En el primer estudio, titulado *¿Pueden obtenerse efectos de congruencia de respuesta en el priming enmascarado de decisión léxica?* (Can Response Congruency Effects Be Obtained in Masked Priming Lexical Decision? Fernández-López et al., 2019, *Journal of Experimental Psychology: Learning, Memory, and Cognition*), nos propusimos comprobar si las personas aplican las mismas instrucciones a los estímulos *prime* y test en la tarea de decisión léxica con *priming* enmascarado, examinando el efecto de congruencia de respuesta con *primes* no relacionados. El efecto de congruencia de respuesta en la decisión léxica puede definirse como la diferencia en los tiempos de respuesta entre un ensayo congruente (es decir, cuando el *prime* obtiene la

misma respuesta que el test; por ejemplo, (árbol-CASAL: "palabra"; geuze-POLCI: "no palabra") y un ensayo incongruente (por ejemplo, geuze-CASAL: "palabra"; árbol-POLCI: "no palabra"). El examen de este efecto nos permite analizar el estado del léxico después de que el estímulo *prime* haya sido procesado en una etapa temprana (Loth & Davis, 2010). En concreto, examinamos si había diferencias en la influencia de un *prime* enmascarado no relacionado que variaba en su estatus de semejanza a la palabra: los primes podían ser una palabra de alta frecuencia (radio), una palabra de baja frecuencia (boina), una pseudopalabra (geuze) o una cadena de consonantes (xfgtp).

Para obtener una imagen completa, diseñamos tres escenarios que variaban en la dificultad de la discriminación de palabras vs. pseudopalabras: pseudopalabras ortográficamente legales emparejadas en elementos subsilábicos con las palabras-test (Experimento 1; por ejemplo, CASAL vs. POLCIR), pseudopalabras ortográficamente ilegales (Experimento 2; por ejemplo, CASAL vs. SYKDD) y pseudopalabras ortográficamente legales sin vecinos (Experimento 3; por ejemplo, CASAL vs. LEDUL). Se observó que, cuando se utilizaban pseudopalabras ortográficamente legales (POLCIR, LEDUL; Experimentos 1 y 3), los tiempos de respuesta a las palabras y a las pseudopalabras no variaban en función de la semejanza de las palabras con los *primes*. Esto sugiere que las personas no pudieron establecer el estatus léxico del *prime* enmascarado. Es decir, no pudieron determinar la probabilidad de que fuera una palabra. Sin embargo, encontramos un escenario diferente cuando se utilizaron pseudopalabras ilegales (SYKDD; Experimento 2): los tiempos de respuesta de las palabras eran más largos cuando el *prime* enmascarado era una cadena consonante que cuando era una pseudopalabra ortográficamente legal (xfgtp-CASAL > geuzel-CASAL). Además, las respuestas a las pseudopalabras fueron más rápidas cuando fueron precedidas por un *prime* compuesto por una cadena consonántica que cuando fueron precedidas por un *prime* ortográficamente legal (xfgtp-SKYDD < geuzel-SKYDD)—pero sólo en los cuantiles iniciales de las distribuciones de tiempo de reacción. Sin embargo,

se podría argumentar que la inclusión de pseudopalabras ilegales en el Experimento 2 cambió la naturaleza de la tarea de una decisión léxica a una tarea de decisión de legalidad ortográfica.

Estos resultados tienen importantes implicaciones para los principales modelos de reconocimiento de palabras (modelo del lector bayesiano; Norris y Kinoshita, 2008; modelos de activación interactiva; Rumelhart, 1981; Grainger y Jacobs, 1996; Coltheart et al., 2001; Davis, 2010). El modelo del Lector Bayesiano (BR) predice tiempos de respuesta similares para las palabras precedidas por un *prime* no-relacionado, independientemente de su semejanza con la palabra en la decisión léxica. En consonancia con el modelo BR, todas las condiciones de *priming* no-relacionado se comportaron de forma similar en los escenarios estándar de reconocimiento de palabras. Sin embargo, cuando la discriminación palabra vs. pseudopalabra era mucho más fácil (Experimento 2; CASAL vs SYKDD), encontramos un efecto de los *primes* no-relacionados. ($\text{xfgtp-CASAL} > \text{geuzel-CASAL}$; $\text{xfgtp-SKYDD} < \text{geuzel-SKYDD}$). No obstante, dado que la tarea en el Experimento 2 puede considerarse una decisión de legalidad ortográfica, el modelo BR puede predecir respuestas "sí" más lentas cuando el *prime* está compuesto por una cadena de consonantes que cuando es un estímulo ortográficamente legal, y viceversa: Como todos los estímulos que no son palabras violan las restricciones ortográficas, todo lo que es legal debe ser una palabra. En cuanto a los modelos de activación interactiva, la falta de influencia del efecto de semejanza a la palabra del *prime* en el procesamiento del estímulo-test cuando se comparan palabras y pseudopalabras legales (Experimentos 1 y 3) ofrece apoyo a la hipótesis de la inhibición lateral selectiva (Davis y Lupker, 2006). Sin embargo, la naturaleza de la tarea en el Experimento 2 (es decir, la decisión de legalidad ortográfica) iría más allá del alcance de la familia de modelos de activación interactiva. Por lo tanto, nuestros hallazgos ofrecen apoyo empírico a la cuenta del lector bayesiano y a la hipótesis de inhibición lateral. Además, se ha demostrado cómo la naturaleza de la decisión requerida por la tarea modula los efectos del *priming* enmascarado.

A pesar de las numerosas ventajas de la técnica de *priming* enmascarado, no es lo suficientemente sensible para detectar manipulaciones sutiles (por ejemplo, luz-luz vs. lz-luz) o examinar los procesos de reconocimiento de letras y palabras en poblaciones especiales, como las personas lectoras en desarrollo o con dislexia. La razón es que, debido a la corta asincronía entre el *prime* y el test, la magnitud de los efectos del *priming* enmascarado es pequeña, en el rango de 10-50 ms en los estudios conductuales. Para superar esta limitación, Lupker y Davis introdujeron la técnica de *sándwich* en 2009. Esta técnica es una variante del *priming* enmascarado estándar en la que la adición de un *pre-prime* (idéntico al estímulo-test) durante 33ms, magnifica el tamaño de los efectos. Sin embargo, los mecanismos específicos en juego, y cómo difieren de la técnica estándar, no han sido completamente dilucidados (véase Lupker & Davis, 2009 frente a Trifonova & Adelman, 2018). En el segundo estudio de esta sección, *Desvelando el impulso en la técnica sándwich de priming* (Unveiling the boost in the sandwich priming technique. Fernández-López et al., 2021, *Quarterly Journal of Experimental Psychology*), nos propusimos caracterizar la técnica del *sándwich* examinando si los efectos del *priming* simplemente se magnifican o, por el contrario, son los procesos subyacentes los que se alteran. Para ello, llevamos a cabo un experimento de decisión léxica con *priming* de identidad comparando directamente las técnicas de *sándwich* y estándar en un diseño intrasujeto. Además, para comprobar la naturaleza de las diferencias entre los métodos estándar y *sándwich*, examinamos el locus del impulso del efecto en la técnica *sándwich* y el curso temporal de los efectos para ambas técnicas mediante gráficos delta y funciones de precisión condicional. Los resultados mostraron que el locus del impulso en la técnica del *sándwich* era doble: respuestas más rápidas en la condición de identidad (a través de un cambio en las distribuciones de tiempos de reacción) y respuestas más lentas en la condición no-relacionada. La similitud de las distribuciones del tiempo de respuesta para los pares de identidad en las técnicas de *sándwich* y estándar—difieren solamente en su localización, pero no en la forma—sugiere que el tiempo extra del *pre-prime* no está alterando los procesos subyacentes más allá de la codificación (véase Gómez et al., 2013, para una discusión de los efectos en las

distribuciones de los tiempos de reacción). Sin embargo, para los *primes* no relacionados, el *pre-prime* en la técnica de *sándwich* dificultó el procesamiento de los pares no relacionados. Este hallazgo advierte de que se debe aplicar cautela a la hora de interpretar algunos de los efectos de *priming* entre condiciones en la técnica de *sándwich*, ya que estos efectos se pueden derivar no de la facilitación en la condición relacionada, sino de la inhibición en la condición no relacionada.

En resumen, los dos estudios presentados aquí ofrecen una valiosa contribución a los aspectos metodológicos de la investigación sobre el reconocimiento de palabras. De ellos se pueden extraer tres aspectos destacados: (i) a la hora de elegir la línea de base, no importa la semejanza de las palabras de los *primes* no relacionados (siempre que la tarea emplee pseudopalabras ortográficamente legales); (ii) la naturaleza de la decisión requerida por la tarea modula los efectos de *priming* enmascarado; (iii) al emplear la técnica del *sándwich*, se debe utilizar tanto las condiciones de identidad como las de no-relación como controles.

We are absurdly accustomed to the miracle of a few written signs being able to contain immortal imagery, involutions of thought, new worlds with live people, speaking, weeping, laughing. We take it for granted so simply that in a sense, by the very act of brutish routine acceptance, we undo the work of the ages, the history of the gradual elaboration of poetical description and construction, from the treeman to Browning, from the caveman to Keats. What if we awake one day, all of us, and find ourselves utterly unable to read? I wish you to gasp not only at what you read but at the miracle of its being readable.

Vladimir Nabokov, *Pale Fire*

Introduction

A book is made from a tree. It is an assemblage of flat, flexible parts (still called “leaves”) imprinted with dark pigmented squiggles. One glance at it and you hear the voice of another person, perhaps someone dead for thousands of years. Across the millennia, the author is speaking, clearly and silently, inside your head, directly to you. Writing is perhaps the greatest of human inventions, binding together people, citizens of distant epochs, who never knew one another. Books break the shackles of time—proof that humans can work magic.

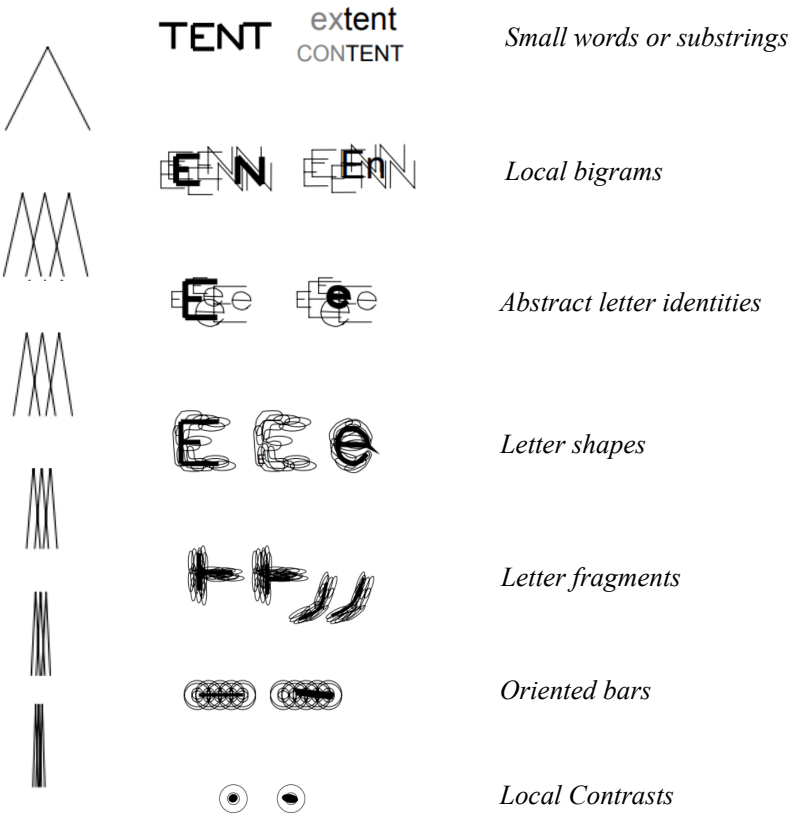
Carl Sagan

Reading is an indispensable skill in our society. In 1859, Abraham Lincoln gave a talk on the most important inventions in human history. He named clothing, stone tools, the wheel, and writing and reading. “Writing—the art of communicating thoughts to the mind, through the eye—is the great invention of the world” (Fehrenbacher, 1989, p.7). Effectively, without reading, we could not access the past, the history, the advances, and the setbacks. It is close to miraculous that Homo sapiens dramatically increased their memory abilities by making some marks on a stone. Importantly, this art of “communicating thoughts through the eye” is one of the most complex and astonishing skills learned by humans. When reading, the eyes fixate most words in the text being read, and typically only once. This implies that readers get a foveal glimpse of most words, and that an essential first stage of skilled reading behavior concerns the processing performed during that glimpse. Thus, although the act of reading entails several important processes, visual word recognition is an integral aspect. Indeed, readers can recognize visually presented words with apparent ease, but substantial machinery lies behind them. When we read a single word, there are a series of complex cognitive processes involving computations at different levels of processing (e.g., perceptual, orthographic, phonological, morphological, syntactic, and semantic) during a few hundred milliseconds.

These processes become evident when a difficulty affects one specific function but not the others. Therefore, it is crucial to investigate in depth each of them, as well as their relationships. Consequently, a large body of research has

tried to shed light on each of these processes, and the present thesis is intended to contribute to the reading research corpus. Specifically, the research conducted in this dissertation focuses on the “cellule” of reading: the letter. If there is one aspect of visual word recognition on which there is a general consensus among researchers, it is the crucial role of individual letters. It is generally assumed that the recognition of a printed word is mediated by the processing involving the letters that make up the word, at least in languages that use alphabetic orthography. According to the Oxford dictionary, a letter is a character that represents one or more of the elementary sounds used in speech and language. Letters are also the first thing we learn from our teachers and the foundation on which our subsequent education is based. So significant is, that leading models of written word recognition give it a special status (e.g., Local Combination Detectors model; Dehaene et al., 2005).

Figure 1
The Local Combination Detectors model



Note. Adapted from Dehaene et al. (2005; Figure 1)

For instance, the Local Combination Detectors model (Dehaene et al., 2005) a hierarchical neural-inspired model that explains word recognition via a series of local combinations of detectors that are sensitive to increasingly larger fragments of words, from shape fragments to complete words. Each neuron is assumed to pool activity from a subset of neurons at the immediately lower level, leading to an increasing complexity, invariance and size of the receptive field at each stage. As can be seen in the Figure 1, at the lower level, combinations of local oriented bars can form local shape fragment detectors. Next, combinations of fragments can then be used to form local shape detectors, which can detect a letter, but only in a given case and shape. At the next step, the *magical moment* occurs: the abstract letter identities can be recognized by pooling activation from populations of shape detectors coding for the different upper and lowercase versions of a letter. The subsequent stage is composed by neurons sensitive to bigrams (i.e., local combinations of letters), which can be combined, conforming the natural next step, forming morphemes or small words. Importantly, words would not be never encoded by a single neuron, but by a sparse distributed population of them.

The general objective of the present dissertation is to examine the processes that occur in the *magical moment*. That is, those processes that build the bridge between the low-level stages of visual processing and the higher-level processing of words: orthographic processing. The orthographic processing refers to the encoding of information about the identities and position of the letters, thus allowing us to differentiate KISS from HISS or GOD from DOG. Therefore, to further our understanding of reading, this dissertation focuses on the small—but crucial—features that the big machinery needs to work, and how the literacy experience shapes them.

To this end, the background and current status of the subject in question will be described in three sections: (i) the emergence of orthographic processing, (ii) the resilience of letter detectors, and (iii) the methodological contributions to visual word recognition research. The section on the emergence of orthographic processing comprises three studies that examined the development of orthographic processing with literacy in the early years of reading acquisition.

Importantly, with reading acquisition running its course, a flexible and powerful decoding system is built: readers can rapidly access the abstract representations of letters despite their format (*a A a A a a a* share the same abstract representation). In the section on the resilience of letter detectors, I present three studies that examined how visual distortion affects orthographic processing (i.e., the mapping from the visual input onto abstract word representations during normal reading). To that end, we focused on a single type of distortion: letter rotation. Finally, the last section includes a methodological contribution to the word recognition research. Specifically, it comprises two studies conducted to further comprehend the cognitive underpinnings of masked priming, the gold-standard tool for studying the very first moments of visual word recognition.

The emergence of orthographic processing

C'est pourquoi nous crions: enseignement! science! Apprendre à lire, c'est allumer du feu; toute syllabe épelée étincelle.

Victor Hugo, *Les Misérables*

Reading acquisition is a major step in child development. If we could look at ourselves when we first started to learn to read, we might remember how we struggled to decode letters and their combinations. How long it took reading a simple word like `radio`, letter by letter, phoneme by phoneme. Nowadays, however, reading is so much part of daily life that we do not realize how complex this skill is, and how hard it was to acquire. Indeed, it takes years of arduous work before the brain engine that supports reading runs so smoothly that we forget it exists.

Learning to read involves a fine-tuning of our primate brain. Using all available resources, our brain recycles the most appropriate regions of our visual cortex for the novel task of word recognition. Nothing in our evolution could have prepared us to absorb language through vision. However, brain imaging shows that the adult brain contains fixed circuits exquisitely adapted to reading (see Dehaene, 2009). Thus, beginning readers draw on established cognitive and sensory abilities that are recruited, modified, and coordinated in novel ways to establish specific strategies of information processing that are optimized for text (Lachman & van Leuven, 2014).

One of the central landmarks of learning to read is the emergence of orthographic processing. This process constitutes the bond between the low-level stages of visual processing and the higher-level processing of words. Specifically, the orthographic processing is responsible for the encoding of letter identity and letter position and order, thus distinguishing between `fat`, `cat`, and `act`. The encoding of letter identities, understood as the mapping of visual input to abstract

letter representations, would occur in the first two years of reading acquisition (Jackson & Coltheart, 2001; Gomez & Perea, 2020). However, it is still not clearly understood when and how the encoding of letter order and position arises. The present section focuses on the emergence of these processes via the examination of location-invariant processing (responsible for encoding the order of letters) and location-specific processing (responsible for encoding the position of letters).²

Importantly, the question of how the brain encodes the order and position of letters within words is a central issue for all leading models of visual word recognition (e.g., Spatial Coding model: Davis, 2010; Overlap model: Gomez et al., 2008; Open Bigram model: Grainger & van Heuven, 2004; Bayesian Reader model: Norris et al., 2010; SERIOL model: Whitney, 2001). Nevertheless, these models focused on an extant perspective (i.e., assuming a fully-developed word recognition system) rather than a developmental perspective. Moreover, despite being critical factors of efficient reading, the empirical evidence on the emergence of location-invariant and location-specific processes is very scarce and inconclusive. Keep in mind that comparing the results from children before and after beginning with reading acquisition may lead to interpretative issues, as the accuracy and latency data vary dramatically across groups (see Perea et al., 2016).

The main aim of the studies presented in this section is to fill this gap by examining the emergence of orthographic processing and the factors that may influence it. Pinpointing the factors that influence the emergence and development of orthographic processing is crucial, as it allows us to identify children in need of pre-literacy intervention and, consequently, to minimize eventual reading difficulties. With this purpose in mind, we conducted the first study of the present section, entitled *Which factors modulate letter position coding in pre-literate children?* (Fernández-López et al., 2021). Our goal was to scrutinize the precursors of letter order coding. Specifically, we examined the

² These two processes are interconnected. It would be difficult to talk about order without considering position and vice versa.

potential links between the encoding of letter order in pre-literate children with the abilities related to early reading acquisition (i.e., basic cognitive processes: attention, sensation, perception, and memory). Importantly, better readers encode letter order more accurately than the worse readers (see Gómez et al., 2021; Pagán et al., 2021). Thus, it is crucial to shed light on the roots of letter-order processing, defining its precursors.

Importantly, research with skilled readers has shown that, once reading is mastered, the rank of the letters changes from being mere visual objects to enjoying a specific status. For instance, when presenting a string of symbols (£\$?%@) individuals fixate in the middle of the string. However, when a string of letters is presented (FGJGM), they fixate also in the exterior letters, especially in the first letter of the string (Tydgat & Grainger, 2009). This phenomenon is due to the location-specific processing, attained by adapting the mechanisms of visual object processing to the constraints of visual word processing. Moreover, the encoding of the order of letters becomes more flexible, leading to the location-invariant processing: FGJM is easily confounded with FJGM, similarly, oragne is easily confounded with orange and cholocate with chocolate. This phenomenon is usually captured when using a same-different matching task in which participants have to decide if pairs of stimuli are the same or different. Results show that “no” responses to a transposed-letter pair (FGJM–FJGM) are slower and more error-prone than the responses to the replacement-letter control (FGJM–FPCM). This effect, known as the transposition letter (TL) effect, also occurs with strings composed of symbols, numbers or unknown letters, but the effect is are substantially larger for strings of letters. The presumed explanation for the dissociation letters vs. other visual objects is that there is an orthographic-specific mechanism used to encode location invariant letter-in-word-order in addition to general positional uncertainty. Crucially, this orthographic-specific mechanism has been posited to emerge with literacy acquisition (Dandurand et al., 2010; Duñabeitia et al., 2014, 2015).

Nevertheless, as indicated above, given the constraints of doing research with children, the empirical data on the emergence of location-invariant

processing is very scarce and inconclusive. One option to overcome this limitation is to design a laboratory analog of learning to read with adult participants. This was the strategy followed in the second and third studies of the present section. In the second study, *Does orthographic processing emerge rapidly after learning a new script?* (Fernández-López et al., 2021), we tested whether orthographic processing, examined via location-invariant and location-specific processing, emerged quickly with literacy. To that end, we designed a study with pre and post-training tasks in which adult readers were trained to read and write in an artificial script during six sessions.

Finally, in the third study, *A letter is a letter and its co-occurrences: Cracking the emergence of location-invariance processing*, we tested whether location-invariant processing emerges rapidly after the exposure to orthographic regularities—bigrams—in a novel script. To that end, we designed a study with two phases. In Phase 1, individuals were exposed to a flow of artificial words for a few minutes, with four bigrams occurring frequently. In Phase 2, participants performed a same-different matching task to examine the potential differences between pairs with a transposition of letters in a frequent (trained) versus infrequent (untrained) bigram.

Which factors modulate letter position coding in pre-literate children?

Abstract

One of the central landmarks of learning to read is the emergence of orthographic processing (i.e., the encoding of letter identity and letter order): it constitutes the necessary link between the low-level stages of visual processing and the higher-level processing of words. Regarding the processing of letter position, many experiments have shown worse performance in various tasks for the transposed-letter pair `jugde-JUDGE` than for the orthographic control `jupte-JUDGE`. Importantly, 4-y.o. pre-literate children also show letter transposition effects in a same-different task: `TZ-ZT` is more error-prone than `TZ-PH` (Perea et al., 2016). Here we examined whether this effect with pre-literate children is related to the cognitive and linguistic skills required to learn to read. Specifically, we examined the relationship of the transposed-letter effects found in the Perea et al. (2016) study with the scores of these children in phonological, alphabetic and metalinguistic awareness, linguistic skills, and basic cognitive processes. To that end, we used a standardized battery to assess the abilities related with early reading acquisition. Results showed that the size of the transposed-letter effect in pre-literate children was strongly associated with the sub-test on basic cognitive processes but not with the other subtests. Importantly, identifying children who may need a pre-literacy intervention is crucial to minimize eventual reading difficulties. We discuss how this marker can be used as a tool to anticipate reading difficulties in beginning readers.

Key words: learning to read, orthographic processing, cognitive processing, pre-literate, transposed-letter effect

Introduction

Whereas language is a unique and sophisticated human ability that emerges naturally in children, reading is a learned skill that needs intensive practice. In fact, reading acquisition is a complex process that involves functional brain changes and requires the correct execution of numerous mental functions (see Maurer et al., 2005). For this reason, children must have adequate perceptual and cognitive skills before the initial steps of reading instruction. Once acquired, reading becomes the most important tool for knowledge acquisition in academic settings and beyond.

In alphabetic scripts, readers can quickly map the visual input onto abstract letter representations and, subsequently, onto word representations (see Dehaene et al., 2005; Grainger et al., 2008). The emergence of these abstract letter representations would occur during the first two years of reading acquisition (Jackson & Coltheart, 2001). Consistent with this view, using Forster and Davis (1984) masked priming technique, Gomez and Perea (2020) found that, for Grade 2 readers, the identification time of a word like `EDGE` is virtually the same when rapidly preceded by the physically identical prime `EDGE` and when preceded by the nominally (but not physically) identical prime `edge`.

Importantly, the process of visual word recognition requires not only the encoding of the abstract identity of the letters that compose each word but also the encoding of the serial order of the words' letters. If this process were absent, we would not be able to distinguish similarly spelled words like `spot` and `stop`. Notably, in a recent paper with adult readers, Schmitt and Lachmann (2020), demonstrated that when a target has to be identified in a string, processing occurs in serial order (i.e., from left to right) for letter stimuli, but not for strings composed of letters from an unknown alphabet (Cyrillic and Hebrew). At the same time, a considerable wealth of experiments with children and adults have shown that the encoding of letter order is only approximate: readers often perceive jumbled words (e.g., `JUGDE` or `CHOLocate`) as the original words (see Castles et al., 2007; Guerrero & Forster, 2008; Lupker et al., 2008; Perea & Lupker, 2003, 2004). As serial order processing is a key component of a wide range of

psychological processes, from perception to action (Logan, 2021), it is not surprising that the encoding of serial order is also an essential part of reading and literacy. The main goal of the present study is to shed some light on which cognitive factors are associated with pre-readers' ability to encode letter position accurately.

In the context of reading development, Castles et al. (2007) proposed a “lexical tuning” model in which children encode progressively more precisely the letter positions within words. For instance, in a series of masked priming experiments, they found that the prime *drak* was much more effective at activating *DARK* in Grade 3 than in Grade 6 children. The rationale of this model is that, as reading abilities develop, letter position coding becomes more accurate (see Perfetti's [2017] lexical quality hypothesis, for a similar claim; but see Grainger & Ziegler [2011], for a different view³). Evidence supporting the lexical tuning model has been obtained not only with children of different ages but also with children of the same age: better readers encode letter position more accurately than the worse readers (see Gómez et al., 2021; Pagán et al., 2021, for evidence with children; see also Andrews & Lo, 2012; Perea et al., 2016, for parallel evidence with adult readers). Furthermore, a poor encoding of serial order may lead to reading difficulties. Friedmann and Gvion (2001) were the first to report that some individuals present problems at encoding letter position, making frequent errors of letter migration within words—reading *broad* for *board*. This deficit, which has been termed “letter position” dyslexia, has been found in many different languages, including English (Kohnen et al., 2012; see Güven & Friedmann, 2019, for a recent review).

Somewhat surprisingly, examining how the encoding of letter order emerges in young readers and whether some pre-existing abilities may help encode serial order in pre-readers has been overlooked in the literature. One of

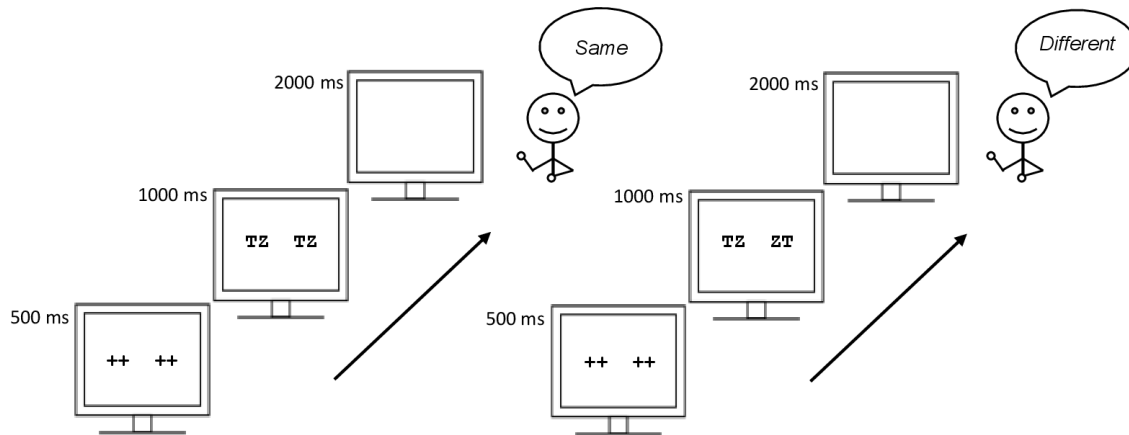
³ In fairness to Grainger and Ziegler (2011), their dual-route model of visual-word recognition focuses on the initial steps of learning to read. The serial letter encoding, which involves precise letter position encoding, would emerge first in reading development (phonological route). Later in development, the parallel encoding of the word's letters (orthographic route) would make letter position coding coarser.

the few exceptions is the longitudinal experiment conducted by Duñabeitia et al. (2015). They used a same-different task with two sequentially presented four-letter strings, and children had to decide whether the letter strings were the same or different. The “different” trials were composed of pairs with two transposed letters (transposed-letter pairs; e.g., *rzsk-rszk*) and pairs with two replaced letters (replacement-letter pairs; e.g., *rzsk-rhck*). If letter position coding is flexible, transposed-letter pairs would be perceptually more similar than replacement-letter pairs, thus producing worse performance (e.g., more false positives). The “transposed-letter” effect is the difference in performance between these two conditions. Duñabeitia et al. (2015) tested the children three times: (1) in their year before pre-school ($M = 4.24$ years), (2) in their pre-school year ($M = 5.21$ years), and (3) in the first year of primary school ($M = 6.32$ years). They only found a transposed-letter effect when the children had learned to read (first-grade children; more error responses for transposed [42.9%] vs. replaced-letter pairs [30.6%]). Duñabeitia et al. (2015) concluded that “position uncertainty emerges as a consequence of literacy training” (p. 549).

An interpretive issue in the Duñabeitia et al. (2015) experiment is that the pre-readers performed very poorly in the same-different task and the sensitivity index, d' , was close to zero for both for replaced and transposed conditions (Perea et al., 2016). This pattern suggests that their version of the same-different task was too difficult for the pre-readers (i.e., the working memory load probably exceeded the children's capacity; see Riggs et al., 2006); thus, one cannot make any inferences on these data. To draw firm conclusions on the encoding of the serial position of letters in pre-readers, Perea et al. (2016) simplified some elements of Duñabeitia et al.'s (2015) same-different task: 1) They used two-letter string pairs instead of four-letter string pairs, 2) the pairs were presented simultaneously instead of sequentially, and 3) the responses were done verbally (i.e., saying “same” vs. “different”) instead of manually (pressing one of two buttons) (see Figure 1).

Figure 1

Depiction of the same-different task in the Perea et al.'s (2016) experiment



Along with “same” pairs (TZ–TZ), Perea et al. (2016) included the following “different” pairs: transposed-letter pairs (TZ–ZT), one-letter replacement pairs (TZ–PZ), and two-letter replacement pairs (TZ–PH). They found a sizeable transposed-letter effect in 4 years-old children (i.e., pre-readers). Specifically, the number of false positives (i.e., “same” responses) was greater to transposed-letter strings (TZ–ZT) than to 1 or 2 replacement-letter strings (TZ–PZ; TZ–PH). Perea et al. (2016) concluded that this pattern reflected a noisy perception of location order, common to all visual objects (see Gomez et al., 2008), rather than an effect that emerges with literacy. Notably, while not analyzed in their paper, shortly after conducting their experiment, Perea et al. (2016) collected the scores of these children in a battery of abilities related to early reading acquisition in Spanish (BIL battery; Sellés et al., 2008).

In the present study, we aim to take a step forward by exploring the potential precursors of letter position coding in pre-readers. To that end, we examined the relationship between the ability of pre-literate children to encode accurately the order of letters—taken from the Perea et al. (2016) experiment—with the five subtests related to reading readiness and subsequent reading success from the BIL battery (Sellés et al. 2008): phonological and alphabetic awareness, metalinguistic knowledge, linguistic skills, and basic cognitive processes. The examination of this issue is important not only at a theoretical

level but also at a practical level. Before learning to read, children must have acquired some perceptual, cognitive, and linguistic skills. Defining the early precursors of precise coding of letter position will shed light on the roots of the processing of serial order when reading letters in words. These analyses would allow us to identify children who may present some deficit (e.g., some mild forms of letter position dyslexia) and start intervening as soon as possible, preventing future reading difficulties.

Thus, in the present study, we examined the relationship between the sensitivity of the readers to distinguish transposed-letter pairs from identity pairs (e.g., TZ-ZT vs TZ-TZ) and the scores of pre-readers (M = 4.5 years old) in phonological and alphabetic awareness, metalinguistic knowledge, linguistic skills, and basic cognitive processes (visual perception and sequential auditory memory) in the BIL battery (Sellés et al. 2008). We focused on transposed letter pairs, as the mechanisms employed to discriminate TZ-ZT from ZT-ZT are based exclusively on letter order. The predictions are clear. In adult readers, basic cognitive processes such as spatial and visual attention have been assumed to play a key role in encoding letter position (see Gomez et al., 2008; McCann et al., 1992). If this generalizes to pre-readers, we expect a positive relationship between the abilities at discriminating TZ-ZT and the scores in these basic cognitive processes. This outcome would imply that educators could use this simple same-different task with a transposed letter pairs to predict reading readiness before starting with the reading instruction. Furthermore, it may also operate as an incentive to design other tasks for pre-readers on perceptive and executive skills to prevent—or at least minimize—potential difficulties at locating letters within words during learning to read. In addition, we expect no relation between linguistic factors (i.e., phonological and alphabetic awareness, metalinguistic knowledge, and linguistic skills) and the sensitivity at distinguishing TZ-ZT in pre-readers—at the time of the experiment, the children did not know the consonant names.

Method

Participants

They were the 20 preschoolers ($M = 4.54$ years; $SD = 3.6$; 7 girls) from a private school of Valencia (Spain). All of them were native speakers of Spanish with no learning developmental problems. An informed consent from their parents was obtained before running the experiment, and the study was approved by the Experimental Research Ethics Committee of the University of Valencia. At the time of testing, the preschoolers were starting to learn the vowels but they did not know the name or sound of the consonant letters (as confirmed by results of the BIL battery).

Procedure

The experiment took place individually in a quiet room within the school premises. DMDX software (Forster & Forster, 2003) was employed for stimulus presentation and recording of the responses. A depiction of the procedure in the same-different task can be found in Figure 1. Accuracy was stressed in the instructions. Ten practice trials preceded the 64 experimental trials. Moreover, the children were assessed with a battery of abilities related to early reading acquisition in Spanish (BIL battery; Sellés et al., 2008).

Materials

For the same-different task, the stimuli were 64 pairs of consonant strings made of two consonants. There were 16 trials in each of the conditions: (1) same pairs (TZ–TZ), (2) transposed letter pairs (TZ–ZT), (2) one-letter replacement pairs (TZ–PZ) and (4) two-letter replacement pairs (TZ–PH). Four counterbalanced lists were created in a Latin square manner, so that each stimulus was rotated across the different conditions. The presentation of the items was randomized for each participant.

To assess the abilities related to early reading acquisition we employed the BIL battery (Sellés et al., 2008). This battery comprises five subtests: phonological awareness, alphabetic awareness, metalinguistic knowledge, linguistic skills, and basic cognitive processes—we obtained a score from each

subtest. For the goals of the present study, we focused on the subtest measuring basic cognitive processes. This subtest explores a series of cognitive processes that take place when we face reading: (1) attention, which leads the mind to concentrate on specific stimuli; (2) sensation (i.e., detection and differentiation of sensory information); (3) perception, which integrates sensory experiences and interprets them for recognition and identification (i.e., giving meaning to what has been selected and picked up at the attentional and sensory level), relying on the patterns stored in the (4) memory. To that end, the subtest assesses the child's sequential auditory memory and the ability to visually discriminate between similar letters and symbols (the child had to circle the symbols that were the same as a target; see Sellés et al., 2008, for a depiction of the other sub-tests).

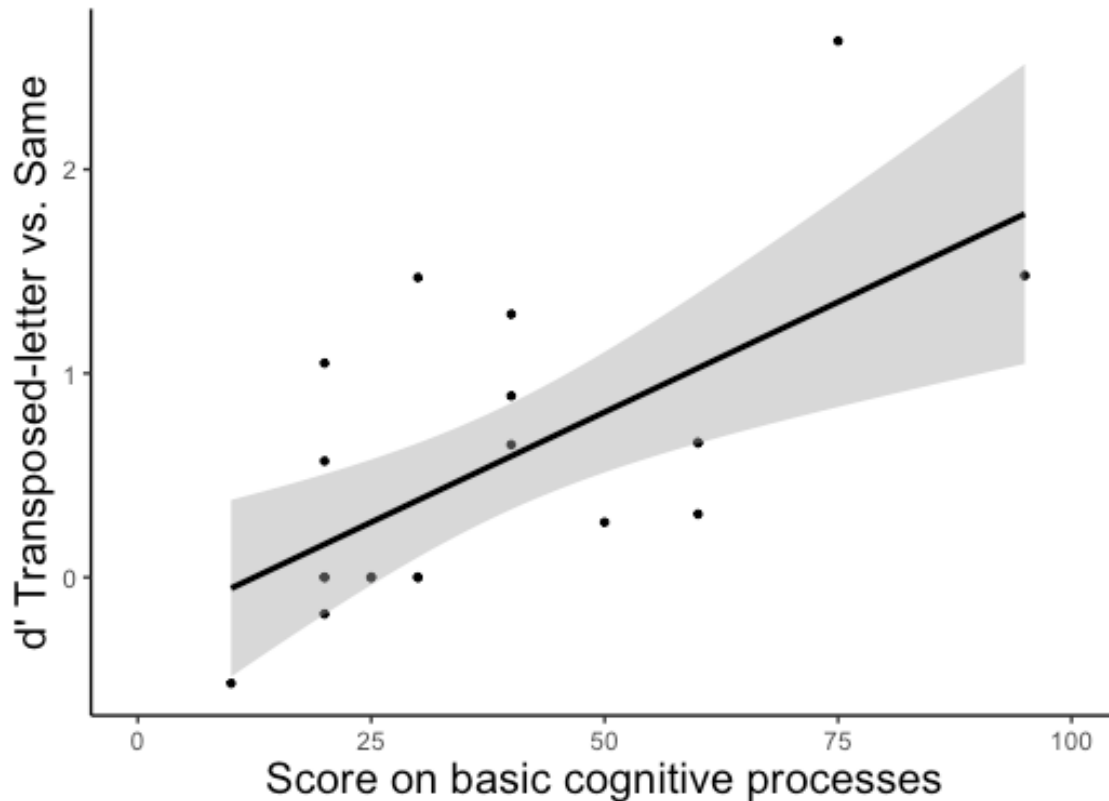
Results

To test whether better pre-reading skills (as measured by the BIL battery) were associated with better performance at differentiating between same and transposed-letter pairs in the same-different task, we conducted frequentist and Bayesian correlation analyses with JASP (Faulkenberry et al., 2020). To compute the Bayes factors, we used the default Cauchy distribution (centered around 0 and with a width parameter $\delta = 0.707$) (see Rouder et al., 2009; Wagenmakers et al., 2017, 2018, for discussion). Specifically, we examined the relation between d' (a measure of sensitivity obtained from the accuracy data of Perea et al., 2016) and the percentile scores in sub-tests of the BIL battery—of note, these findings were virtually the same if we had employed the raw scores from the sub-scales. For the computation of d' , we used the hit rate for same trials and the false alarm rate for the transposed letter trials (TZ-TZ vs. TZ-ZT)—in signal detection theory, chance-level performance [$d' = 0$ or no sensitivity] occurs when the hit rate for the identical items is the equal to the false alarm rate for the different items. Of note, mean accuracy for same trials in the Perea et al. (2016) was .83; for different trials, it was .33 for transposed-letter strings and .68 and .88 for one-letter and two-letter replacement strings, respectively.

Results of the correlational analyses in the present study showed that those children who better differ transposed-letter from “same” pairs ($TZ-ZT$ vs. $TZ-TZ$) had the higher scores in the sub-test on basic cognitive processes ($r = .634$, $p = .003$; see Figure 2). Indeed, the alternative hypothesis was 18.6 ($BF_{10} = 18.559$) times more likely than the null hypothesis with the present data (see Jeffreys, 1961, for interpretation of Bayes Factors). In addition, there were no signs of a relationship between the children’s performance differentiating between same and transposed letter pairs and the other (linguistic) subtests (all $ps > .24$, $BF_{s10} < 0.528$).

Figure 2

Correlation between the d' when discriminating transposed letter pairs from same pairs and the scores on the basic cognitive processes of BIL



For completeness, we explored the relationship between performance in the replacement letter conditions and the BIL battery; to this end, we computed separately d 's for same vs. one-letter replacement trials ($TZ-TZ$ vs. $TZ-PZ$) and for same vs. two-letter replacement trials ($TZ-TZ$ vs. $TZ-PH$), and then calculated the correlations between these two d 's and the sub-test of the BIL battery. None of

these correlations produced evidence in favor of a relationship (all $ps > .147$; all $BF_{10} < 0.478$).

Discussion

Identifying the cognitive precursors of reading is vital to determine those children who are ready to start learning to read and those who still need some cognitive maturation or some early intervention. This would prevent later reading difficulties and disorders and the frustration and psychological discomfort that such problems usually entail. With this matter in mind, in the present study, we scrutinized the roots of the mechanisms underlying the encoding of letter position in strings (i.e., one of the critical factors of efficient reading; see Castles et al., 2007; Logan, 2021). Specifically, we examined the relationship between the capability of pre-literates to differentiate between transposed-letter pairs and identity pairs (e.g., TZ-ZT vs. TZ-TZ) and these children's scores in basic cognitive processes. Results showed that the pre-literate children who best differentiated between TZ-ZT and TZ-TZ in a same-different task were those with higher scores on basic cognitive processes (see Figure 2). Notably, the subtest of basic cognitive processes was not generically associated with sensitivity in the same-different task (i.e., it was not related to performance for replacement-letter trials); instead, it is uniquely associated with accuracy in letter position coding. Thus, at the theoretical level, this outcome reflects that basic cognitive skills shape the ability to encode serial order in letter strings (e.g., a smaller value of the parameter responsible for perceptual uncertainty in models of letter position coding; see Davis, 2010; Gomez et al., 2008). Furthermore, at an applied/educational level, our findings imply that a simple same-different task can be used to assess reading readiness: the better the performance in this task, the better the encoding of letter order, diminishing the chances of letter position dyslexia.

In addition, our findings suggest that the preparing-to-reading arise early in development with some non-specialized processes that would be recruited and adjusted to guide the subsequent functional reading progress (see Lachmann &

van Leeuwen, 2014). Further support to this idea can be found in the study of Saygin et al. (2016). They found that the cortical location of the visual word form area (i.e., the brain region specialized for letter string; Dehaene & Cohen, 2011) at age 8 (when children read) can be predicted by the distinctive connectivity of the same region at age 5 (pre-literates). Taken together, these studies emphasize that early detection of deficiencies in the visual analysis of the input is crucial to prevent later reading difficulties (Friedmann & Gvion, 2001; Shetreet & Friedmann, 2011). This is consistent with the assumption that children with reading difficulties have a general impairment in domains other than linguistic (e.g., an impairment in multisensory integration; see Gori & Facoetti, 2014; Lachmann & van Leeuwen, 2014). Therefore, there are possibly many (complementary) ways to test whether pre-readers are prepared to starting reading learning (e.g., the same-different task or the “avatar task”; see Perea et al., 2014); this would be a valuable endeavor for future studies.

We acknowledge that the present study comes with some limitations. Firstly, because of the correction criteria of the BIL test, it was not possible to obtain separate scores for the tasks that make up the basic cognitive processes sub-test (sequential auditory memory and visual discrimination). Furthermore, although the processes assessed in the basic cognitive sub-test were not linguistic in nature, the stimuli contained symbols, letters and words, thus, making it difficult to clearly disentangle basic cognitive processes and verbal processes. Future tests should be more specific to characterize all possible aspects that shape the cognitive processes of pre-literates. In addition, it would have been desirable to have obtained further data from the same children once they started reading learning. These data would have allowed us to test whether the findings in the same-different task with transposed letters in pre-literate children were a good predictor of letter position coding once children acquired knowledge about letters. Furthermore, these longitudinal data would have also allowed us to examine the interplay between the emergence of orthographic processing during learning and the scores in cognitive and linguistic processes. Indeed, once the children start to read, other elements would begin playing a significant role, such as alphabetic knowledge or phonologic awareness (see Dehaene et al., 2015).

A complementary strategy for future research would be to run parallel longitudinal same-different experiments on serial order using to-be-learned letters vs. unknown letters (e.g., letters from another alphabet). The data pattern should be similar for the pre-readers for both types of stimuli, but one would expect differences when the children learn to read. Critically, these differences could be considered as markers of the emergence of orthographic processing (see Grainger, 2018). While this approach is ideal on an a priori basis, it suffers from various methodological issues. One would need to design a feasible task for children of different tasks that minimizes both ground and ceiling effects. However, it is challenging to create a task achievable for pre-literates and complicated enough to draw differences among developing readers. For instance, deciding whether two 4-letter strings are the same is extremely challenging for pre-literates, whereas deciding whether 2-letter strings are the same may be too easy for developing readers (see Perea et al., 2016, for discussion). As a result, it is very difficult to experimentally study the emergence and development of orthographic processes in pre-readers. To further complicate matters, there are also other potential limitations, such as the lack of control for prior letter knowledge and other linguistic elements in pre-readers, or that the duration of experiments for pre-readers would need to be quite short to keep them attentive. An alternative is to design laboratory analogs of children's reading acquisition that consists of training adults to read a novel script (see Fernández-López et al., 2021; see also Chetail, 2017; Taylor et al., 2011). This approximation is not as ecological as one would desire (see Maurer et al. 2010; Taylor et al., 2017), but it definitively increases the control on the process of acquiring the novel orthography.

In sum, the early identification of potential problems that may slow down reading development is of fundamental importance for psychologists and educators. In the present study, we found that those pre-readers who performed better in basic cognitive processes tended to be those who would encode more accurately letter position in a simple same-different task. This finding highlights that learning to read should not be based solely on letter knowledge and phonological decoding. We also need to consider that learning to read is built on

a basic cognitive foundation, probably related to multisensory integration based on visual attention.

Does orthographic processing emerge rapidly after learning a new script?

Abstract

Orthographic processing is characterized by location-invariant and location-specific processing (Grainger, 2018): i) strings of letters are more vulnerable to transposition effects than the strings of symbols in same-different tasks (location-invariant processing); and ii) strings of letters, but not strings of symbols, show an initial position advantage in target-in-string identification tasks (location-specific processing). To examine the emergence of these two markers of orthographic processing, we conducted a same-different task and a target-in-string identification task with two unfamiliar scripts (pre-training experiments). Across six training sessions, participants learned to fluently read and write one of these scripts. The post-training experiments were parallel to the pre-training experiments. Results showed that the magnitude of the transposed-letter effect in the same-different task and the serial function in the target-in-string identification tasks were remarkably similar for the trained and untrained scripts. Thus, location-invariant and location-specific processing do not emerge rapidly after learning a new script; instead, they may require thorough experience with specific orthographic structures.

Key words: orthographic processing; training; letter position coding; first-letter advantage; artificial script

Introduction

Reading is an acquired skill that involves some functional brain changes and requires, in alphabetic scripts, associating the letters that compose each word with their appropriate speech sounds. A common assumption in the literature is that, for a mature word recognition system, the process of identifying words comprises a series of stages that map the visual input onto abstract letter representations and, subsequently, onto whole-word representations (see Dehaene et al., 2005; Grainger, 2008; but see Price & Devlin, 2011, for an alternative account). The processing of orthographic representations connects the low-level stages of visual processing to the higher-level linguistic processing of words. These orthographic representations contain information about the identity and order of each of the word's constituent letters, thus allowing readers to distinguish similarly spelled words like `cure` and `core`, which differ in the identity of just one letter, or words like `slat` and `salt`, which differ in the order of two of the letters (see Grainger, 2018, for review). Indeed, the question of how the brain encodes the identity and order of the letters that constitute each word is a central issue for all leading models of visual word recognition (e.g., Spatial Coding model: Davis, 2010; Overlap model: Gomez et al., 2008; Open Bigram model: Grainger & Van Heuven, 2004; Bayesian Reader model: Norris et al., 2010; SERIOL model: Whitney, 2001).

In the present experiments, we examined the emergence of two fundamental markers of orthographic processing after learning a new script: 1) location-invariant processing; and 2) location-specific processing (see Grainger, 2018). For simplicity, the above-cited models focus on an extant perspective (i.e., they assume a fully-developed word recognition system frozen in time) rather than on a developmental perspective. (We defer a discussion of the research focused on the development rather than the emergence of orthographic representations [e.g., Castles et al., 2007; Grainger et al., 2012; Marinus et al., 2018] until the Discussion section.) We first describe in some depth how location-invariant and location-specific processing differ between letters and other visual objects (e.g., symbols, unknown letters). Then, we describe how acquiring a new

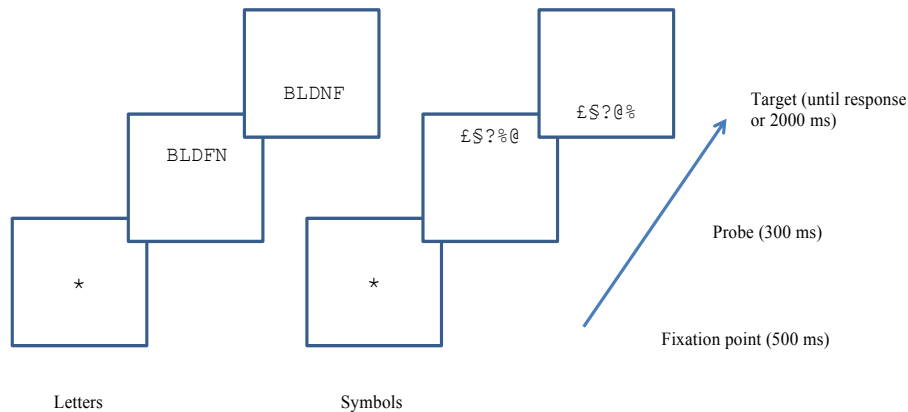
script may affect both phenomena. Finally, we offer a rationale for the two experiments proposed in the current paper.

Location-invariant processing refers to the mechanisms responsible for encoding the “relative positions of a set of object identities” (Grainger, 2018, p. 345) (i.e., the encoding of the order of visual objects [letters] in a string composed of several objects [a word]). This has often been examined with the same-different matching task (see Krueger, 1978; Ratcliff, 1981, for early evidence), as it allows researchers to compare the processing of letters vs. the processing of other types of visual objects. In this task, participants have to decide if two strings of characters presented subsequently are the same or not (see Figure 1). The most studied phenomenon of location-invariant processing is the transposed-letter effect (henceforth, TL effect; see Grainger, 2018, for review). The TL effect refers to the insensitivity of readers to the position of letters compared to the identity of the same letters: “no” responses to the transposed-letter pair FGJM-FJGM (the underline is to emphasize the manipulation) in a same-different matching task are slower and more error prone than the responses to the replacement-letter control FGJM-FPCM. These effects also occur with strings composed of symbols or unknown letters (e.g., £\$?@-£?@ is slower and more error prone than £\$?@-£#<@), which suggests that there is some positional noise in the representations of visual objects in a string (see Gomez et al., 2008; Norris & Kinoshita, 2012). But the critical finding is that transposition effects are substantially larger for strings of letters than for strings of other visual objects (e.g., numbers, symbols, pseudoletters; Duñabeitia et al., 2012; Massol et al., 2013; see also García-Orza et al., 2010; Muñoz et al., 2012). To explain the greater transposition effect for strings of letters than for strings of other visual stimuli, Grainger (2018; see also Marcet et al., 2019; Massol et al., 2013) suggested that, on top of positional noise, there is an orthographic-specific mechanism used to encode location invariant letter-in-word-order. Critically, this orthographic-specific mechanism has been posited to emerge with literacy acquisition (Dandurand et al., 2010; Duñabeitia et al., 2014, 2015). Therefore, when learning a new script, the emergence of location-invariant orthographic

processes would produce an increase of letter transposition effects in a same-different task.

Figure 1

Depiction of the same-different task

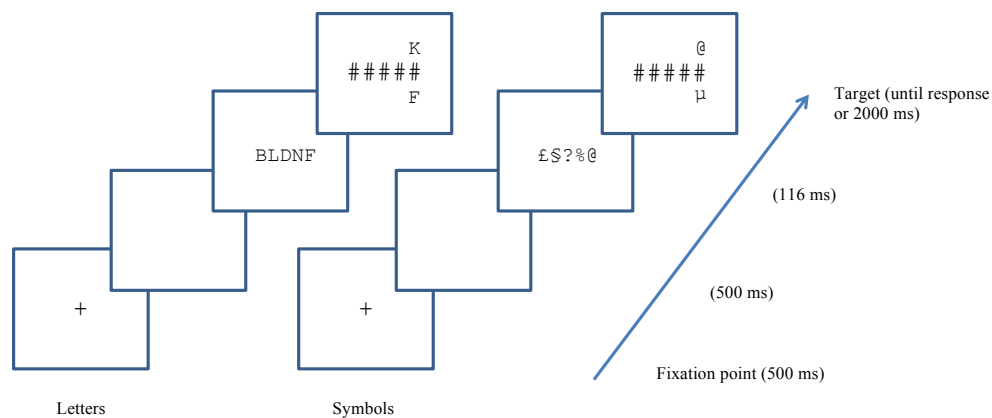


Location-specific processing of visual information refers to the parallel processing of the position of characters (e.g., letters) within one object (e.g., a word). This type of processing has usually been examined with a post-cued partial report target-in-string identification task (henceforth, TSI task), based on the two-alternative-forced-choice task first introduced by Reicher (1969) and Wheeler (1970). In the typical setup of the TSI task applied to location-specific processing (e.g., Tydgat & Grainger, 2009), a string of five characters (e.g., letters: FGJGM; symbols: £\$?%@) is presented briefly while the participant is looking at the middle of the string. The string is subsequently followed by a pattern-mask with a cue indicating one of the positions in the string (see Figure 2 for illustration). The participants' task is to choose, from the two alternatives, the one that matches the identity of the character at the cued location. When presented with strings of symbols, adult readers typically show a Λ -shape function (i.e., an advantage of the middle, fixated position) (Tydgat & Grainger, 2009; see also Chanceaux & Grainger, 2012; Chanceaux et al., 2014; Grainger et al., 2010; Scaltritti et al., 2018; Vejnović & Zdravković, 2015). In contrast, for letter strings, adult readers typically show a W-shape serial position function of accuracy (see Tydgat & Grainger, 2009). That is, there is an advantage not only for the fixated, middle letter, but also of the exterior letters—primarily the initial

letter. The dissociation in serial position function for strings of symbols vs. letters has also been obtained with developing readers (see Ziegler et al., 2010). To explain this pattern, Tydgat and Grainger (2009) proposed the Modified Receptive Field (MRF) theory. The idea is that, at the onset of learning to read, the status of letters changes from being independent visual objects to becoming parts of a higher-order object (i.e., the string). This is attained by adapting the mechanisms of visual object processing to the constraints of visual word processing (see Grainger, 2018; Grainger & Hannagan, 2014; see also Dehaene et al., 2005, for neural correlates). Specifically, the MRF theory assumes that, with reading acquisition, location-specific letter detectors are developed and their receptive fields become reduced in size and elongated to the left—note that the initial letter is critical to translating orthographic representations into phonological representations (Grainger et al., 2016). Thus, the emergence of location-specific processing when learning a new script is expected to produce an initial position advantage in a TSI task.

Figure 2

Depiction of the target-in-string identification task



The empirical data on the emergence of location-invariant and/or location-specific processing are very scarce (see Duñabeitia et al., 2015, for an exception). Duñabeitia et al. (2015) examined the emergence of location-invariant processing in a longitudinal same-different experiment with children from four years (i.e., preliterate children) to six years (i.e., children who had acquired orthographic representations). In the accuracy data, Duñabeitia et al.

(2015) found a significant TL effect for the older children, but not for the preliterate children, and claimed that the “the skills related to the processing of internal characters’ identities and positions are inherently dependent on literacy” (p. 548). However, d' (i.e., a measure of sensitivity) did not differ from zero in the experiment with preliterate children (i.e., children were performing at chance level), which raises questions about any interpretation of the data (see Perea et al., 2016, for discussion)⁴.

The main goal of the present experiments was to overcome this gap by examining whether acquiring a new script affects location-invariant processing and location-specific processing by using a same-different task and a TSI task, respectively. We designed a laboratory analogue of children’s reading acquisition in which adults were trained to read and write in a new, unfamiliar script. As Chetail (2017) indicated, the use of artificial scripts (i.e., sets of characters either from unknown scripts or newly devised) provides a unique opportunity to “examine the developmental course of a given orthographic process which is stable in adults” (p. 103). Furthermore, recruiting adults as subjects allows us to control the participants’ prior knowledge (i.e., we can make sure that participants are not familiarized with the characters; see Maurer et al., 2010; Taylor et al., 2017), and it also enables us to increase the number of conditions and trials of the experiments: adult participants can carry out longer experimental sessions than children, and this allows us to ensure appropriate reliability. Additionally, comparing the results of preliterate children and developing readers may lead to intricate interpretative issues, as the accuracy and latency data vary dramatically across groups (see Perea et al., 2016). Critically, the behavioral effects of learning an unfamiliar script in adults can be generalized to the effects elicited on children when learning their first language (see Taylor et al., 2011, for discussion).

⁴ While the Perea et al. (2016) experiment with preliterate children rules out an interpretation of TL effects as being fully dependent on literacy acquisition, it does not provide any insights as to the emergence of location-invariant processing. To test whether location-invariant processing is influenced by literacy in young children, one would need to run a re-test—ideally with letters vs. symbols (or letters from a new alphabet)—after these children learn to read.

In the current study, we employed a classic design with a pre-training phase and a post-training phase. The pre-training phase comprises two experiments: a same-different experiment on the TL effect (i.e., testing location-invariant processing; Experiment 1) and an experiment with a TSI task on the serial position function (i.e., testing location-specific processing). The pre-training experiments were conducted using eighteen consonant letters from an artificial monospaced font (BACS2serif; Vidal & Chetail, 2017). This font was used to create two different scripts: Script 1 (11 letters; two vowels and nine consonants) and Script 2 (11 letters; two vowels and nine consonants) (see Figure 3). One of the scripts was learned via print-to-sound training along five days and the other was used as a control. An important issue is the choice of the appropriate control script. Keep in mind that the letters in the trained script would not only activate print-to-sound correspondences, but they would also be visually familiar. That is, after training, the pseudoletter “ $\mathfrak{d}$ ” would not only correspond to a phoneme, but it also would become a familiar object. Thus, one could argue that any effects from the trained script in the post-training phase could be merely due to visual familiarity. To separate the effects of learning to read from the effects of visual familiarity, participants were familiarized with the visual form of the characters of the control script.

In the pre-training phase of Experiment 1, we conducted a same-different task using five-letter strings. In the pre-training phase of Experiment 2, we conducted a TSI task with five-letter strings. In both experiments, Script 1 and Script 2 were presented in separate blocks. Subsequently, each participant received training in one of the scripts (trained script) and was familiarized with the characters of the other script (control script). For the print-to-sound training, each individual learned the grapheme-phoneme associations of nine consonant letters and two vowel letters from one of the two scripts across six training sessions: half of the participants learned the letters in Script 1 and the other half learned the letters in Script 2. Prior research has shown that readers can show some expertise in a new script quite rapidly. For instance, Chetail (2017) found that individuals acquired new regularities (e.g., letter and bigram frequency effects) after a relatively short amount of time, even in unfamiliar and complex

scripts. Likewise, Brem et al. (2018) reported that two hours of training were enough for individuals to show some expertise for a novel script (e.g., an increase of the N1 amplitude). For the visual familiarization with the control script, participants were exposed to the eleven characters of the script, but without mentioning any orthographic or phonological information.

Figure 3
Association between the letters of the new scripts and their corresponding phonemes

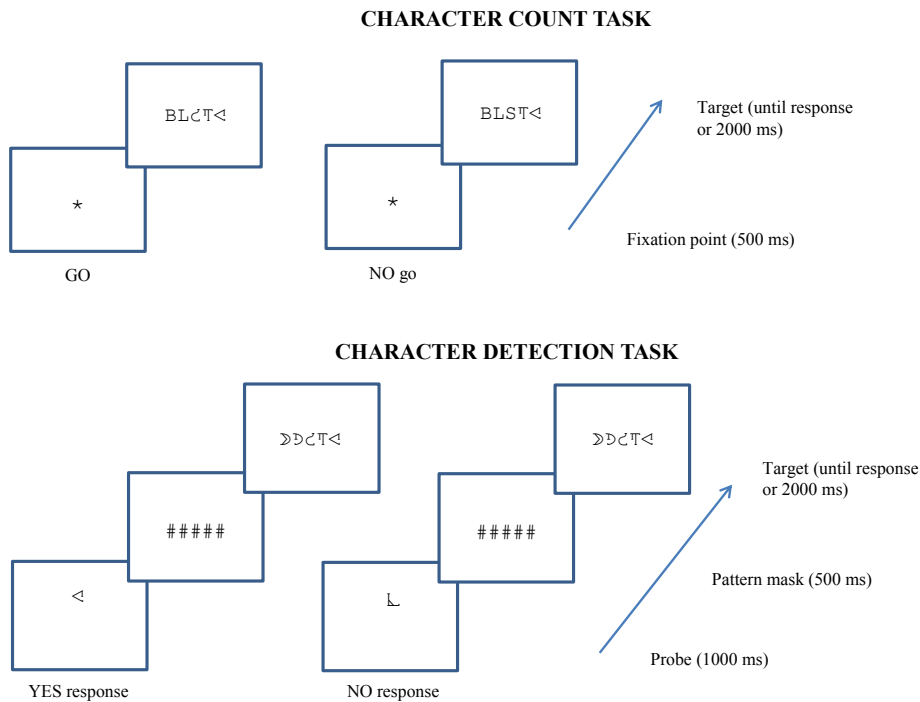
Script 1	Script 2	Phoneme
A	⊖	/a/
∖	⌘	/i/
⊗	⋈	/d/
≡	┌	/k/
⊕	└	/f/
⋈	⊕	/l/
∪	π	/p/
τ	⊖	/m/
⊥	?	/r/
⋈	⋈	/s/
⊗	⋈	/θ/

On day 1, participants first learned the grapheme-phoneme associations in the trained script. As in Spanish—their native language, all the grapheme-phoneme associations are transparent (e.g., the letter “i” always corresponds to the phoneme /i/). Importantly, the print-to-sound training allows us to ensure that participants will learn the new script as a group of letters and not as mere symbols or shapes. As Chetail (2017) pointed out, “a critical feature that distinguished letters from other symbols or shapes is that letters are used to transcribe speech according to a structure code” (p. 110). To consolidate the learning of the trained

script over the next five sessions, which took place in a window of five working days, participants were asked to read aloud and write down series of items of increasing length, from 4-letter to 8-letter strings. Furthermore, the participants were familiarized with the visual forms of the characters of the control script. To that end, on each training session, each individual performed a character detection task and a character count task (see Figure 4 for depiction of each task; see also Chetail, 2017, for a similar strategy).

Figure 4

Depiction of the character count task (top panel) and the character detection task (low panel)



Finally, all participants had to pass a final test on the sixth day to show that they successfully acquired the experimental script. Once the individuals had passed this test (the criterion was set at 20 or more correct responses out of 24 in reading/writing within a time limit), they took part in the post-training experiments. What we should note here is that the print-to-sound training occurred in absence of semantics. Recent research has emphasized the role of

sound-based strategies when learning to read a new script (e.g., see Brem et al., 2018; Taylor et al., 2017, for evidence with adults). This print-to-sound training also enabled us to isolate the orthographic processes from semantic processes, thus minimizing any influences from top-down processes.

The post-training experiments were parallel to the experiments from the pre-training phase. They were designed to test whether location-invariant and location-specific processing has emerged in the trained script—for comparison purposes, we also conducted a block with stimuli in an overlearned script (i.e., Roman alphabet). The predictions were clear. If location-specific orthographic processes emerge after literacy acquisition—as proposed by Dandurand et al. (2010) and Duñabeitia et al. (2014, 2015), we would expect a greater TL effect in a same-different task for the trained script in the post-training phase than in the pre-training phase. Keep in mind that, in the post-training phase, the TL effect for the newly learned script would have two constituents: a positional noise component—shared with the pre-training phase—and an orthographically-based component. For the control script, the TL effect should remain similar in magnitude in the pre- and post-training phases: the TL effect would be due to perceptual uncertainty. Alternatively, if TL effects were similar in magnitude in the pre- and post-training phase for the two scripts, this would reveal that the emergence of location-invariant processing does not emerge quickly after learning print-to-sound correspondences in a new script.

Second, if location-specific letter detectors are formed with literacy acquisition—as proposed by the MRF theory (Tydgate & Grainger, 2009), the trained script would elicit a first letter advantage in the TSI task when measuring the serial position function in the post-training phase. This would be accompanied by an advantage of the fixed, middle position i.e., a Λ -shape function) for the two scripts in the pre-training phase and for the control script in the post-training phase. Alternatively, if the pattern of data still shows a Λ -shape function for both the trained and untrained scripts in the post-training phase, this would suggest that print-to-sound training does not rapidly lead to the emergence of location-specific orthographic processing.

Experiment 1: location-invariant processing

Method

Participants

The sample was composed of twenty-eight university students, all of them native speakers of Spanish with normal/corrected-to-normal vision and with no history of reading or hearing disorders. All of them signed an informed consent form before participating in the experiment. Participants received a small monetary compensation after the experiment. In consonance with the registered protocol, the final number of participants was determined via a Sequential Bayes Factor design maximal n (see Schönbrodt & Wagenmakers, 2018, for the advantages of this approach) starting with a sample size of 28 participants. To compute the Bayes factors for the critical interaction (i.e., the three-way interaction between Phase x Script x Probe-target relationship in the accuracy data) required for the sampling procedure, we obtained the Bayes factors in the by-subjects Bayesian ANOVA—note that all the stimuli were strings of random consonants (or consonants from artificial scripts) so generalization over participants was more important than generalization over random consonants. This Bayes factor was computed in JASP (Wagenmakers et al., 2018) as the ratio of the model that contained the factor of interest (i.e., all the main effects, two-way interactions, and the three-way interaction) vs. the model that did not contain the effects of interest (i.e., all the main effects and the two-way interactions). This Bayes Factor exceeded 6 (i.e., the criteria established in the pre-registered protocol) in favor of the null hypothesis ($BF_{10} = .081 \rightarrow BF_{01} > 12$), so sampling was stopped with $n = 28$.

Materials

We created 240 five-consonant string pairs (probe and target) in Script 1 (see Figure 3), 240 five-consonant string pairs in Script 2 (see Figure 3), and 240 five-consonant string pairs in Roman alphabet (using Courier New font; e.g., STNGB). The two artificial scripts stemmed from the same font: BACS2serif (Vidal & Chetail, 2017). The string pairs in Script 1 and Script 2 were presented in

separate, counterbalanced blocks—the string pairs in the Roman script were presented at the end of the post-training phase. All character strings were composed of non-repeated letters. There were 120 “different” pairs and 120 “same” pairs for each character string type. For the “different” pairs in each script, 60 pairs were created by transposing two adjacent letters (e.g., $\mathcal{C}\mathcal{T}\mathcal{L}\mathcal{D}\mathcal{V}$ - $\mathcal{C}\mathcal{L}\mathcal{T}\mathcal{D}\mathcal{V}$; 1st-2nd, 2nd-3rd, 3rd-4th, 4th-5th), and 60 pairs were created by replacing two adjacent letters (e.g., $\mathcal{C}\mathcal{T}\mathcal{L}\mathcal{D}\mathcal{V}$ - $\mathcal{C}\mathcal{S}\mathcal{Q}\mathcal{D}\mathcal{V}$; 1st-2nd, 2nd-3rd, 3rd-4th, 4th-5th). Thus, each block contained 240 pairs of character strings. In total, each participant was given 480 trials in the pre-training phase (240 string pairs in Script 1 and 240 string pairs in Script 2) and 720 trials in the post-training phase (240 string pairs in Script 1, 240 string pairs in Script 2, and 240 string pairs in Roman script). The proportion of transpositions/replacements was the same in all possible locations. To counterbalance the probe-target pairs, we created two lists for each script (see Massol et al., 2013, for a similar procedure). For the practice phase, we also created eight five-consonant string pairs for each block (eight pairs in the Script 1 block, eight pairs in the Script 2 block, and eight pairs in the Roman alphabet block).

For the learning to read sessions, we created a template in a standard presentation software with the graphemes of the new script and the associated phoneme (see Figure 3), which were recorded by a female voice and digitalized at a sampling rate of 44.1 kHz. For each script, we created 18 items and 18 utterances of 4-characters and 5-characters, respectively, 30 items and 30 utterances of 6-characters, 66 items and 66 utterances of 7-characters, and 120 items and 120 utterances of 8-characters. The items and utterances were grouped separately by lists of 12 items that were presented with standard presentation software.

To have the participants familiarized with the characters of the control script, we employed a character detection task and a character count task (see Figure 4). For the character detection task, we created a total of 144 pairs of items in each script (i.e., probe and target). The probe was always a single character and the target consisted of a string of artificial characters with different length (18 items of four and five characters, respectively, 30 items of 6-

characters, 66 items of 7-characters and 120 items of 8-characters.). For the character count task, we employed, for each script, 18 items of four and 5-characters, respectively, 30 items of 6-characters, 66 items of 7-characters and 120 items of 8-characters. Each string could be composed by artificial characters or by artificial and Roman characters.

Procedure

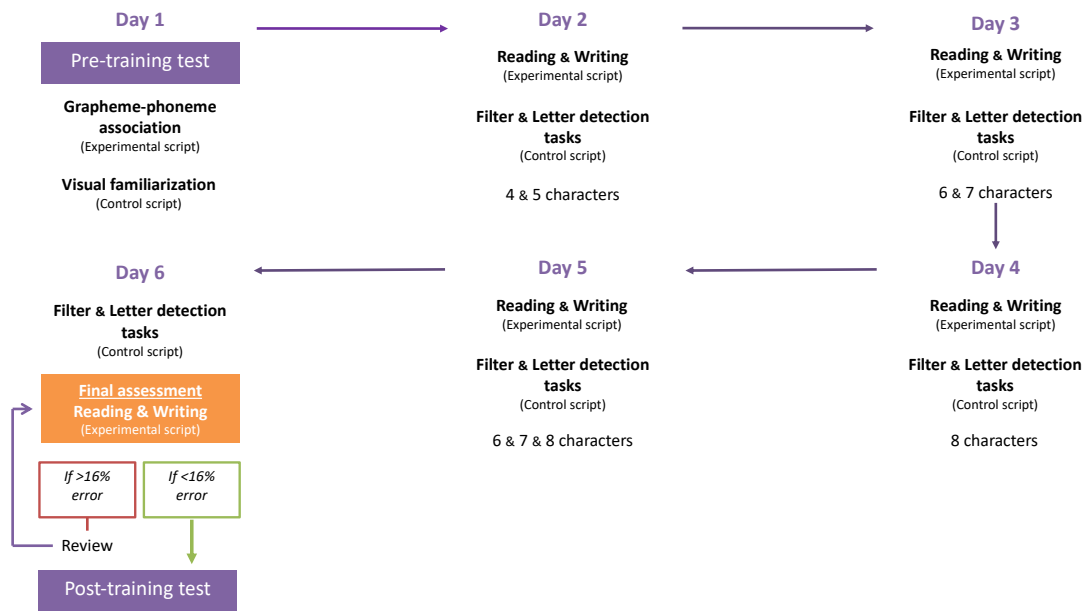
Pre-training-test. Participants were tested either individually or in groups of two in a quiet room. DMDX software (Forster & Forster, 2003) was used to display the sequence of stimuli and to register the timing/accuracy of the responses. Response times were measured from target onset until the participant's response. On each trial, a fixation point (*) was displayed for 500 ms in the center of a computer screen. Next, the fixation point was replaced by a probe, which was presented for 300 ms and positioned 3 mm above the center of the screen. Then, the target item appeared one line 3 mm below the center of the screen. The target remained on the screen until the response or 2,000 ms had passed. All stimuli were presented in a monospaced font (15 pt BACS2serif for Scripts 1 and 2; 15 pt Courier New for the Roman letters) in black on a white background. Participants were told that they would be presented with two strings of consonants and that they would have to press the "yes" key if they were the same, and they were asked to press the "no" button if they were different (see Figure 1). Participants were instructed to make this decision as quickly and as accurately as possible. Eight practice trials preceded the 240 experimental trials in each block. Participants did not receive feedback during the experiment. The session lasted for 18-22 minutes.

Training. Participants were trained individually in the presence of the experimenter along a window of six working days in a quiet room (see Figure 5). Half of the participants learned the grapheme-phoneme associations in Script 1 and the other half learned the grapheme-phoneme associations in Script 2. There were two blocks in each session: for the trained script, participants received print-to-sound training (i.e., grapheme-phoneme association) and, for the control

script, they participated in tasks that entailed visual familiarization with the control characters—the order of each block was counterbalanced across sessions.

Figure 5

Schematic depiction of the training sessions



On the first day, after the pre-training experiments, participants learned the association between the spoken forms of each grapheme in one of the unknown scripts (i.e., the experimental script: Script 1 or Script 2; see Figure 3) and they also familiarized with the visual form of the control script (i.e., the script not used for the grapheme-phoneme association). For the print-to-sound training, the characters were presented on a computer screen with their corresponding sound (participants could click on the character with the mouse and listen to its associated phoneme). Participants were also asked to hand-copy the new letters on a sheet of paper and they were given as much time as they needed to learn these associations.

For the visual familiarization part, the characters of the control script were presented one by one on a computer screen in absence of any orthographic or phonological information. Participants had to hand-copy the control characters on a sheet of paper, without a time deadline. Then, all the characters of the

control script were presented and participants took as much time as they wanted to familiarize with them (see Chetail, 2017, for a similar procedure). The experimenter checked and corrected (when necessary) the writing of the letters and characters to minimize the differences in handwriting quality⁵.

On the second day, participants were presented with items of four and five characters; on the third day, they were presented with items of six and seven characters; on the fourth day, they practiced with items of eight characters and, on the fifth day, participants were presented with items of six, seven and eight characters (see Figure 5). For the print-to-sound training, the general procedure was as follows. Participants had to read aloud and write down 36 items without time deadline. The items were presented on the computer screen, divided into alternating blocks (reading and writing) of 12 items. For reading aloud, a list of 12 items was presented on the screen (e.g., “ㄱㅏㅑㅑ”) and participants were asked to read the items one by one (e.g., /daki/). During this task, the experimenter provided feedback after each item (i.e., correct/incorrect response). If the participant made a mistake, the experimenter encouraged her/him to read it again on her/his own. If the participant could not figure out the correct response, the experimenter indicated it, remarking the grapheme-phoneme correspondences. For the writing blocks, a list of 12 “loudspeaker” signs (🔊) were presented on the screen. Participants were asked to press the 🔊 sign, listen to the pronunciation (e.g., /daki/) and then write down the corresponding graphemes in a sheet of paper (e.g., “ㄱㅏㅑㅑ”). As blocks of 12 items were presented simultaneously, participants were able to see and listen to each item as many times as needed. The experimenter provided feedback after each item (i.e., correct/incorrect response). If the participant made a mistake, the experimenter asked her/him to listen the item again. If the participant could not figure out the mistake on her/his own, the experimenter told her/him the correct response remarking the grapheme-phoneme correspondences.

⁵ Exact handwritten copies of the character were not required, as neither are exact copies of the letters when children learn to write. It was enough if the handwritten copy of the character approximated to the original to be identified and distinguishable from the others characters.

For each block in both tasks (i.e., reading aloud and write down), the experimenter annotated the mistakes (if any), the order of the tasks, performing times and other comments (e.g., the most repeated errors) in an assessment form. Moreover, after each block of 12 items, the experimenter provided general feedback of performance (i.e., correct responses, type of errors, and timing). Importantly, on the first sessions (days 2 and 3), the learning goal was to correctly establish the grapheme-phoneme correspondences. Thus, the experimenter focused mainly on the errors made by the participants. Then, on sessions 4 and 5, when the participants hardly made any mistakes, the experimenter encouraged them to read and write down as fast as possible.

For the visual familiarization block with the control script, participants completed a character detection task and a character count task in each training session. The items had the same length as the items of its corresponding print-to-sound training session. For both tasks, DMDX software (Forster & Forster, 2003) was used to display the sequence of stimuli and to register the timing/accuracy of the responses. On each trial of the character detection task (see Figure 4), a probe was presented for 1,000 ms one line 3 mm above the centre of the screen. The probe was subsequently replaced by a pattern mask with same length as the subsequent target (“#####”) on the centre of the screen for 500 ms. Then, the target appeared and remained on the screen until response or 2000 ms had passed. All stimuli were presented in a monospaced font (15 pt BACS2serif) in black on a white background. Participants were told that they would be presented with a character and then with a string of characters (both in the control script) and they would have to press the “yes” key if the probe appeared in the subsequent string, and they were asked to press the “no” button if the single-character did not appear in the string.

On each trial of the character count task (see Figure 4), a fixation point (*) was displayed for 500 ms on the centre of a computer screen. Next, the fixation point was replaced by the target (i.e., a character string). The target remained on the screen until the response or 2,000 ms had passed. All stimuli were presented in a monospaced font (15 pt BACS2serif and 15 pt Courier New) in black on a white background. Participants were asked to press the “yes” key only when the

item presented was composed of 3 or more characters of the control script—keep in mind that the target items consisted of only control script characters (BACS2serif) or a mixture of characters of the control and Roman scripts. 30% of the items consisted of only one or two control script characters and Roman characters (i.e., participants should not press the “yes” key). Participants received feedback on the general accuracy after each task.

Finally, on the sixth day, before conducting the post-training experiments, participants received a final test with 24 items of 8-characters (12 for reading aloud and 12 for writing) presented in the same format as in the training. They had to do the test in less than 1 min and 30 s, and 3 min and 30 s, respectively.⁶ Those participants who passed the assessment with at least 84% of accuracy (20 out of 24 correct responses) took part in the post-training experiments.⁷

Post-training test. The post-training test was the same as in Experiment 1a, with a final additional block with Roman letters. The post-training tests lasted for approximately 25-30 minutes.

Results

All participants were able to write and read the newly learned script fluently and passed the final training test at the first attempt (see Appendix B, for performance of participants along the training sessions). In accordance with the pre-registered protocol, one participant was replaced because of an overall accuracy level below .60 in the replacement-letter condition.

Different trials

Confirmatory analysis. The dependent variables were response time and accuracy. Error and extremely short responses (less than 250 ms: 0 responses)

⁶ The time limit was set by averaging the reaction times of two pilot participants and adding 30 s more—keep in mind that the pilot participants were members of the lab and they were highly motivated.

⁷ A minimum of 70% accuracy was required in the visual control tasks. All participants met this criterion.

were omitted from the latency analyses—there were no responses longer than the 2-sec deadline (i.e., they were automatically categorized as errors). The mean RTs for the correct responses and the accuracy in each experimental condition are presented in Table 1 (see also Figures 6 and 7). We performed the statistical inference not only using (generalized) linear mixed effects models, but we also computed Bayes Factors.

Table 1.

Mean correct response times (in ms) and accuracy (in brackets) in the different conditions of Experiment 1

	Untrained Script			Trained Script		
	Different		Same	Different		Same
	Transposed	Replaced		Transposed	Replaced	
Pre-training	624 (.550)	602 (.714)	550 (.908)	639 (.551)	624 (.733)	581 (.868)
Post-training	585 (.603)	566 (.783)	540 (.888)	600 (.603)	572 (.795)	538 (.910)

Note. For the Roman script, the correct response times and accuracy (in brackets) were 590 ms (.529) for transposed pairs, 574 ms (.812) for replaced pairs and 511 ms (.931) for same pairs.

In the inferential analyses, we focused on “different” trials, as these are the ones with the TL manipulation. The main research question in the experiment was whether location-invariant processing—as measured by the TL effect—emerges in the trained but not in the untrained script in the post-training phase. To test this hypothesis, we employed (generalized) linear mixed effects (LME) models in R (R Core Team, 2019) using the *lme4* 1.1-21 package (Bates et al., 2019) and the *BayesFactor* 0.9.12-4.2 package (Morey & Rouder, 2018) with three fixed factors: Phase (pre- vs. post-training), Script (trained vs. untrained), and Probe-target relationship (transposed, replaced). Regarding the LME analyses, because of the normality assumption required, the raw RTs were inverse-transformed ($-1000/\text{RT}$). The most complex fitted model that converged was:

$-1000/\text{RT} \sim \text{script} * \text{condition} * \text{phase} + (1 + \text{script} + \text{phase} | \text{subject}) + (1 | \text{item})$

For the generalized linear mixed analyses of the accuracy data, responses were coded as binary values (1 = correct, 0 = incorrect) and we used the `glmer` function in the *lme4* package (family = binomial). The most complex fitted model that converged was:

`accuracy ~ script*condition*phase + (1+script|subject) + (1|item)`

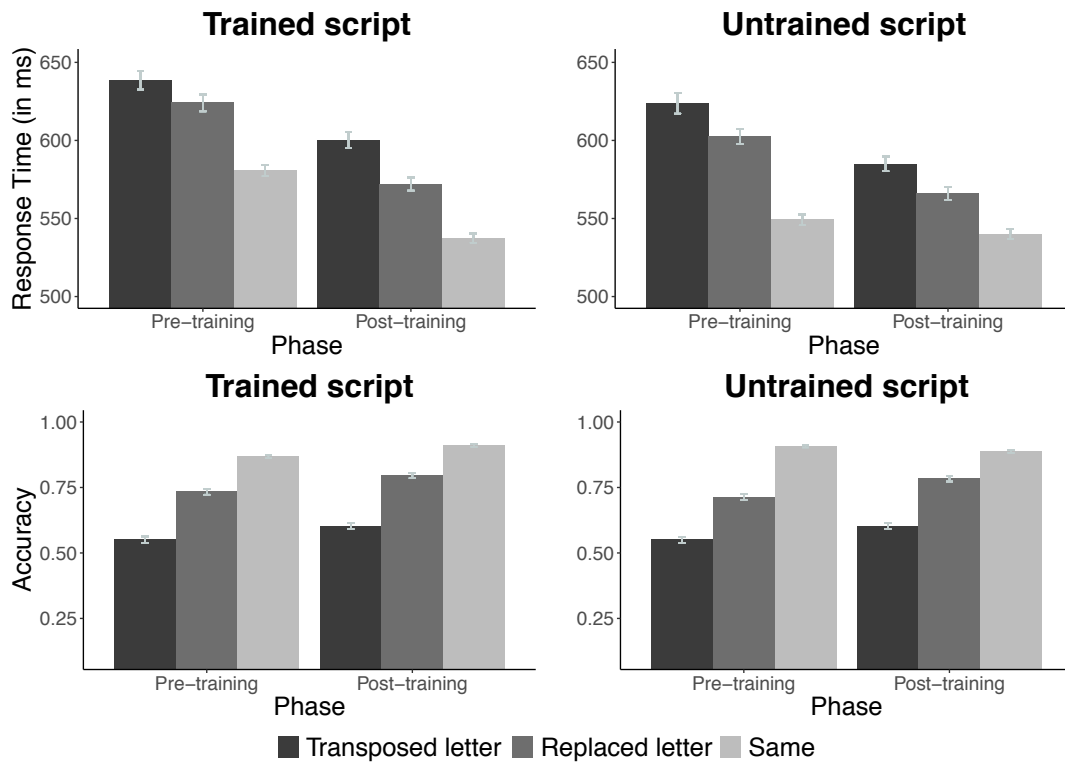
(In Appendix A, we report the [non-pre-registered] analyses using Bayesian linear mixed effects models with the maximal random structure—the results were essentially the same as those reported here.) To compute the Bayes factors on the latency data, we used `lmBF` function from the *BayesFactor* package with the default Cauchy distribution (centered around 0 and with a width parameter $\delta = 0.707$) (see Rouder et al., 2009; Wagenmakers et al., 2017, for discussion). To compute the Bayes factors on the accuracy data, we calculated the Bayes factors from Bayesian analyses of variance (ANOVAs) with the aggregated data by participants—note that items were strings of letters in an artificial script. For the computation of the Bayes factors for each effect, we followed the same logic as in prior research (see Leininger et al., 2017; Staub & Goddard, 2019, for illustration). For the numerator, we compared the maximal model that included the effect of interest vs. a null model that does not assume any fixed effects or interactions. For the denominator, we compared the maximal model after excluding the effect of interest vs. a null model that does not assume any fixed effects or interactions. The ratio between these two Bayes factors was the Bayes Factor of the effect.

Latency analyses. Responses were, on average, 21 ms faster for the RL condition than for the TL condition (i.e., overall TL-effect; 591 vs 612 ms; $b = .08$; $t = 6.18$ $p < .001$; $BF_{10} = 6.69e+05$) and participants responded, on average, 41 ms faster in the post-training tests than in the pre-training tests (581 vs 622 ms; $b = .13$; $t = 4.55$ $p < .001$; $BF_{10} = 20.34$). We found no overall differences in response times between the trained and untrained script ($b = -.01$; $t = -.60$, $p = .55$; $BF_{10} = 1.48$). The interaction between Phase and Script barely reached the significance level in the frequentist analyses ($b = -.04$; $t = -2.43$, $p = .02$), but the Bayes Factors indicated anecdotal evidence towards a null effect ($BF_{10} = 0.55$).

None of the other interactions approached significance (all $t_s < .78$, $p_s > .59$; all $BF_{10} < .35$) (see Figure 6).

Figure 6

Mean reaction times (top panel), accuracies (bottom panel) and standard errors in the trained and untrained scripts of Experiment 1.



Accuracy analyses. Accuracy was higher for the RL condition than for the TL condition (.756 vs .577; $b = -1.05$; $z = -12.73$, $p < .001$; $BF_{10} = 1.071e+39$) and participants were more accurate in the post-training phase than in the pre-training phase (.696 vs .637; $b = -.38$; $z = -4.44$, $p < .001$; $BF_{10} = 415911$). Neither the effect of Script ($b = -.08$; $z = -.83$, $p = .41$; $BF_{10} = .21$) nor any of the interactions approached significance (all $z_s < 1.19$, $p_s > .23$; all $BF_{10} < .26$) (see Figure 6).

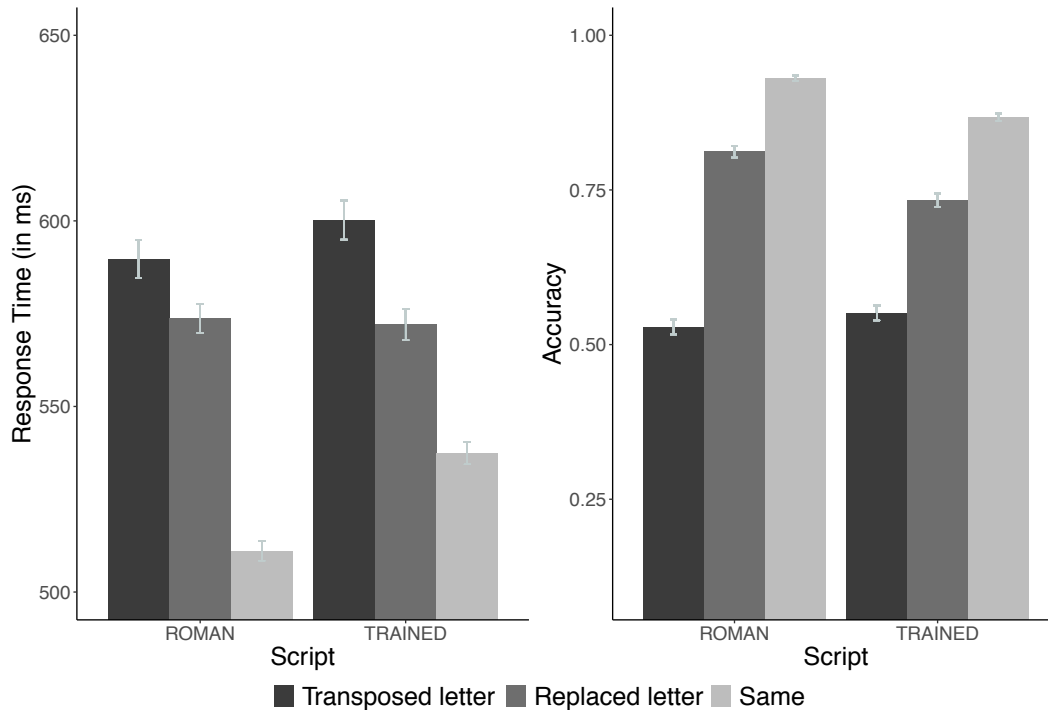
Exploratory analyses. As indicated in the pre-registration, we also compared the TL effect of the trained script (post-training phase) and the TL effect in an overlearned script (i.e., the Roman script). The two fixed factors in the analyses were Script (Trained [post-training] vs. Roman) and Probe-target relationship

(transposed, replaced)—the inferential analyses were parallel to those described above (see Figure 7). The most complex fitted model that converged was:

Dependent_Variable [-1000/RT or accuracy] ~ script*condition + (1+script|subject) + (1|item)

Figure 7

Mean reaction times (left panel), accuracies (right panel) and standard errors in the trained (post-training phase) and Roman scripts of Experiment 1



Latency analyses. Participants responded, on average, 22 ms faster in the RL condition than in the TL condition (573 vs 595; $b = .077$ $t = 6.00$, $p < .001$; $BF_{10} = 129.80$). There were no overall differences between the trained script and the Roman script ($b = .02$ $t = .38$, $p = .71$; $BF_{10} = 0.53$). The interaction between probe-target relationship and script was not significant ($b = -.01$; $t = -.41$, $p = .68$; $BF_{10} = 0.54$).

Accuracy analyses. Participants were more accurate in the RL condition than in the TL condition (.804 vs .566; $b = -1.12$; $z = -13.06$, $p < .001$; $BF_{10} = 3.308e + 22$), whereas the effect of Script was not significant ($b = .22$; $z = 1.14$, $p = .25$; $BF_{10} = 0.734$). We found a significant interaction between the two factors ($b = -.58$; $z = -$

4.67, $p < .001$; $BF_{10} = 7.758$), which reflects that the TL effect was greater in the Roman script than in the trained script (.283 vs .192).

Same trials

While not indicated in the pre-registered protocol, the examination of “same” responses in the pre- and post-test phases for the trained and untrained scripts may shed some light on the role of orthographic-phonological training when processing letter strings. To analyze “same” responses, we employed (generalized) linear mixed effects models on the latency and accuracy data. The two fixed effects were Script (trained vs. untrained) and Phase (pre-training, post-training). The most complex model that converged in the (generalized) linear mixed effects models was:

DV [-1000/RT or accuracy] ~ script * phase + (1 + phase|subject) + (1|item)

These analyses were complemented with Bayesian linear mixed effects models using the maximal random factor structure (see Appendix A).

Latency analyses. Responses were, on average, 26 ms faster in the post-training phase than in the pre-training phase (539 vs. 565 ms; $b = .14$; $t = 2.93$, $p = .01$), whereas there were no signs of an effect of script ($b = .01$; $t = -.51$, $p = .61$). More important, the interaction between Script and Phase was significant ($b = -.12$; $t = -8.09$, $p < .001$). This reflected that responses were faster in the post-training than in the pre-training phase for the trained script (41 ms; 537 vs. 581 ms), but not for the untrained script (9 ms; 540 vs. 549 ms) (see Figure 6)

Accuracy analyses. Participants were more accurate in the post-training than in the pre-training phase (i.e., main effect of phase; $b = -.48$; $z = -2.97$, $p = .003$), and with the untrained than with the trained script (i.e., main effect of script; $b = -.27$; $z = -3.17$, $p = .001$)—note that the effect of script was .009 and was not corroborated by the Bayesian linear mixed effects analyses (see Table A3). More important, mimicking the latency analyses, we found an interaction between the

two factors ($b = .72$; $z = 6.16$, $p < .001$): participants were more accurate in the post-training test than in the pre-training test for the trained script (.910 vs .868, respectively), but not for the untrained script (.888 in the post-training vs .908 in the pre-training test) (see Figure 6).

Discussion

As usual, we found a substantial transposed-letter effect for “different” responses in the same-different task: participants’ responses were faster and more accurate for replacement-letter pairs than for transposed-letter pairs (see Krueger, 1978; Ratcliff, 1981, for early evidence; see also Duñabeitia et al., 2012; Massol et al., 2013; Perea et al., 2016). More important for the purposes of the experiment, the magnitude of the transposed-letter effect was similar for the trained script and for the (visual) control script in both latency and accuracy data. We did find that the responses to “different” trials were, on average, faster and more accurate in the post-training phase than in the pre-training phase. However, this occurred similarly in both the trained and visual control scripts, hence, it could have been due to the participants’ being more visually familiar with the new letters. In addition, the size of the transposed-letter effect was greater in the Roman script than in the newly learned script in the accuracy analyses (28.3% vs. 19.2% of errors, respectively)—note that previous studies on the transposed letter effect also showed significant effects on accuracy, but not on response latencies (e.g., Massol et al., 2013; Perea et al., 2016; Perea & Lupker, 2004).⁸ This is consistent with the idea of orthographic location-invariant mechanisms being at work in the Roman script, but not in the newly learned script (see Duñabeitia et al., 2012; Massol, et al., 2013; see also García-Orza et al., 2010;

⁸ Furthermore, for the Roman script, we found that size of the transposed-letter effect was considerably smaller for external than for internal transpositions: 14.5% vs. 42.1%, respectively (e.g., see Gomez et al., 2008, for a similar pattern). In contrast, for the trained script, the size of the transposed-letter effect was only slightly lower for external transpositions than for internal transposition (17.0% vs. 21.4%, respectively). This again suggests that the transposed-letter effect in the Roman script and the trained script reflect different underlying processes.

Muñoz et al., 2012, for greater transposed-letter effect for letters than for other visual objects [symbols, digits, false fonts]).

The above results may offer the impression that training a new script did not create any stable orthographic representations. However, this interpretation is difficult to reconcile with the fact that “same” responses were substantially faster and more accurate in the post-training phase than in the pre-training phase for the trained script (538 vs. 581 ms; .910 vs .868), but not for the untrained script (540 vs 549 ms; .888 vs .908). This finding strongly suggests that learning to read in the new script helped encoding the letter strings, thus reflecting the emergence of rudimentary orthographic representations.

In sum, the current same-different experiment favors the view that location-invariant processing, as measured by the transposed-letter effect, does not emerge rapidly after learning to read in a new script. We defer a more detailed discussion of this issue in the General Discussion.

Experiment 2: location-specific processing

Method

Participants

They were the same as in Experiment 1. To compute the Bayes factors for the critical interaction (i.e., the three-way interaction between Phase x Script x Position in the accuracy data) required for the sampling procedure, we obtained the Bayes factors in the by-subjects Bayesian ANOVA. This Bayes Factor exceeded 6 (i.e., the criteria established in the pre-registered protocol), $BF_{10} = .05 \rightarrow BF_{01} = 20$, so sampling was stopped with $n = 28$.

Materials

Based on the design and procedure used by Tydgat and Grainger (2009), we created a set of 180 five-consonant strings in two unfamiliar scripts: 90 in Script 1 and 90 in Script 2, and—for the post-training phase—a set of 90 five-consonant

strings in Roman alphabet (e.g., STNGB). None of the letter strings contained repeated characters.

We designed three different blocks (one for script: Script 1, Script 2, and Roman script) with 90 experimental and 9 practice trials each one. The order of the artificial script blocks was counterbalanced between subjects. We manipulated the target position in the array (1st, 2nd, 3rd, 4th, and 5th position). As in the Tydgat and Grainger (2009) experiments, each of the target characters was presented 2 times at each of the five target positions (once above and once below the backward mask), and 40 times at a non-target position (i.e., each target character played as alternative at each of the five positions). Importantly, the incorrect alternative was never presented in the stimulus array. We created two lists for each of the artificial scripts, manipulating the orientation of the target character (i.e., in List 1, the target was presented above the array, whereas in List 2 the same target was presented below the array). These two lists were presented to all participants, one for the pre-training test and the other for the post-training test. The sub-experiment with the Roman script was presented at the end the post-training test.

The learning sessions materials were the same as in Experiment 1.

Procedure

Pre-training test. Participants were tested individually or in groups of two in a quiet room. DMDX software (Forster & Forster, 2003) was used to display the sequence of stimuli and to record the timing and accuracy of the responses. Each trial began with a fixation point (i.e., “+”) that stayed on the screen for 500 ms and was followed by a 500-ms with the blank screen. Then, a string of five letters was presented for 116 ms (see Scaltritti et al., 2018, for the same setup). The array of characters was followed by a backward mask (“#####”) accompanied by two characters, above and below the mask, respectively, at one of the five possible array positions (i.e., characters position as a post-cue) (see Figure 2). The stimuli were displayed on the screen until the participant responded or 2 seconds had passed. All stimuli were presented in black on a white background. We employed a monospaced font for the two scripts. Participants were asked to decide which

of the two characters was present in the corresponding position of the preceding array. They were required to press the “up arrow” key on the keyboard for the character above and the “down arrow” key for the character below the array. They were explicitly instructed to fixate at the center of the array and make the decision as quickly and as accurately as possible. The two scripts were presented in separated blocks, counterbalanced by subjects. Nine practice trials preceded the 90 experimental trials in each of the experimental conditions (90 trials in Script 1 and 90 trials in Script 2). Participants did not receive feedback during the experiment. There was a short break between blocks. The session lasted for around 20-25 minutes.

Training. It was the same as in Experiment 1 (see Figure 5).

Post-training-test. It was the same as in the pre-training test, except for the addition of a third block with stimuli in the Roman script. The session lasted for around 30 minutes.

Results

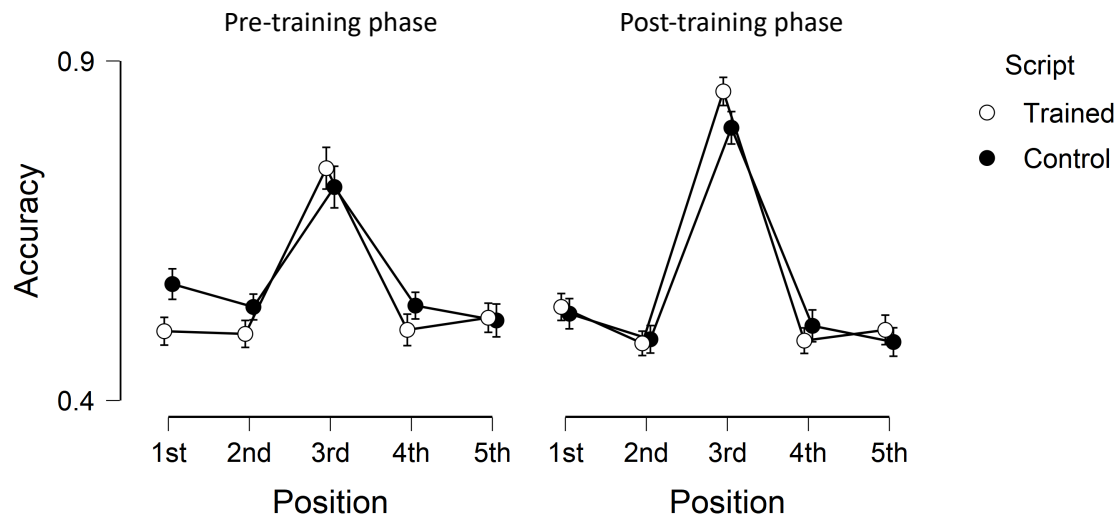
As indicated in the pre-registration, those participants with less than .60 of accuracy in the middle position in the pre- and post-training experiments were replaced—this occurred with four participants. Mean accuracies (and standard errors) for all target types and target positions are presented in Figure 8. The three fixed factors were Phase (pre- vs. post-training), Script (trained vs. control), and Position (1st, 2nd, 3rd, 4th, 5th). By-subjects and by-items classical and Bayesian ANOVAs were performed on the accuracy data. The computation of the Bayes factors was parallel to that described for accuracy in Experiment 1. In Appendix A, we report supplementary [non-pre-registered] analyses using Bayesian linear mixed effects models.

Confirmatory analyses. The ANOVAs showed that accuracy was a function of serial position, $F_1(4, 108) = 67.87$, $p < .001$, $BF_{10} = 2.497e+66$; $F_2(4, 175) = 106.45$, $p < .001$, $BF_{10} = 9.445e+36$. Accuracy levels were higher in the central, third position (.778) than in the first, second, fourth and fifth positions

(with mean accuracy levels of .535, .502, .510 and .507, respectively). Furthermore, the overall accuracy levels were virtually the same for the pre- and post-training phase, $F_1(1,27) = 1.18$, $p = .29$, $BF_{10} = .10$; $F_2 < .001$, $BF_{10} = .09$, and for the trained and control scripts (both $F_s < 1$; both $BF_{s10} < .11$) (see Figure 8).

Figure 8

Mean accuracies and standard errors in the trained and untrained scripts of Experiment 2



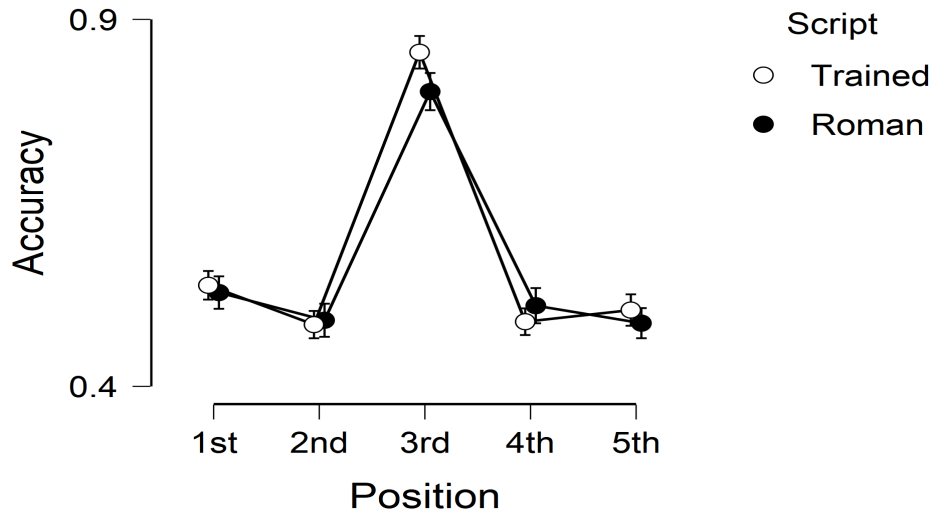
The serial position function differed in the pre-training and post-training phase (Position x Phase interaction; $F_1(4, 108) = 8.65$, $p < .001$, $BF_{10} = 241.21$; $F_2(4, 175) = 6.22$, $p < .001$, $BF_{10} = 53.42$): this occurred because accuracy in the third position was higher in the post-training phase than in the pre-training phase (.828 vs .728, respectively). Neither the interaction between Position and Script ($F_1(4, 108) = 2.65$, $p = .04$, $BF_{10} = .23$; $F_2 < 1$, $BF_{10} = .01$) nor the interaction between Script and Phase ($F_1(1, 108) = 3.22$, $p = .08$, $BF_{10} = .50$; $F_2 < 1$, $BF_{10} = .12$) were significant. Finally, there were no signs of a Phase x Script x Position interaction (both $F_s < 1$; $BF_{s10} < .09$).

Exploratory analyses. We also compared the serial position function of the trained script (in the post-training phase) and an overlearned script (i.e., Roman script) (see Figure 9). The two fixed factors in the ANOVAs were Script (Trained

[post-training] vs. Roman) and Position (1st, 2nd, 3rd, 4th, 5th). The analyses were parallel to those described above.

Figure 9

Mean accuracies and standard errors in the trained (post-training phase) and Roman scripts of Experiment 2



Accuracy was a function of serial position, $F_1(4, 108) = 154.29, p < .001$, $BF_{10} = 4.104e+48$; $F_2(4, 260) = 60.82, p < .001$, $BF_{10} = 7.918e+23$. Participants were substantially more accurate on the third position (.841) than on the other letter positions (.540, .504, .525 and .527, in the first, second, fourth and fifth positions, respectively). In addition, we did not find any clear signs of a difference in the overall accuracy levels in the trained script and the Roman script (.574 vs. .601; $F_1(1, 27) = 3.14, p = .09$; $BF_{10} = .17$; $F_2 < 1$; $BF_{10} = .78$). Finally, as can be seen in Figure 9, the serial position functions of the Roman and trained scripts were remarkably similar and the interaction between the two factors was not significant, $F_1(4, 108) = 2.28, p = .09$; $BF_{10} = .09$; $F_2(4, 260) = 2.05, p = .09$, $BF_{10} = .31$).

Discussion

The current experiment, using a target-in-string identification task, showed an advantage of the middle, fixated position over the other positions for the trained and control scripts not only in the pre-training phase, but also in the post-training phase. We did find a small numerical advantage of the initial position over the second position (see Figure 8), but this difference was similar for the trained and untrained script in the pre-/post-training phases—this pattern was also corroborated in the analyses with Bayesian linear mixed effects models (see Appendix A). We also found that accuracy was higher in the post-training than in the pre-training session—this occurred mainly in the central, fixated letter. This effect was similar for the trained script and for the visual control script, thereby it can be parsimoniously explained in terms of visual familiarity. Taken together, these findings strongly suggest that location-specific letter detectors do not emerge rapidly after learning to read and write in an artificial script.

Finally, in the exploratory analyses, we failed to find a stronger initial letter advantage in the Roman script when compared to the trained script. We prefer to keep cautious about this latter finding. First, in the instructions, we emphasized that participants should be looking at the fixation point at the beginning of each trial. Second, because it was an exploratory analysis, participants always performed the task with Roman letters in the last block. As a result, the attentional capture at the middle position that occurred in initial blocks with the artificial scripts could have been dragged into the last block. A more comprehensive explanation is presented in the General Discussion section below.

General Discussion

We designed two experiments to track the emergence of early orthographic processes in the first stages of learning to read through the examination of two markers of orthographic processing (Grainger, 2018): location-invariant processing (Experiment 1) and location-specific processing (Experiment 2). To that aim, we employed a design with a pre-training phase and a post-training

phase in which adults were trained in reading and writing nonsense words in an artificial script along six sessions. Participants successfully mastered the trained script after the learning-to-read training (see Appendix B). All of them passed the final assessment in the prescribed time with no errors in the reading aloud and writing down tasks. To control for visual familiarity, participants were also familiarized with the visual form of the characters of the control script during the training sessions.

The emergence of location-invariant processing

The first research question was whether readers show some location-invariant processing for the newly learned script on top of the position uncertainty that may affect all visual objects in a string. As stated in the Introduction, location-invariant processing has been proposed to emerge with literacy acquisition (Dandurand et al., 2010; Duñabeitia et al., 2014, 2015).

To our knowledge, the only published study that directly examined this issue was conducted by Duñabeitia et al. (2015). They employed a longitudinal design using a same-different experiment in which a group of children was tested in their antepenultimate pre-school year (i.e., preliterate children; mean age = 4.24 years), in their last pre-school year (i.e., preliterate children; mean age = 5.21 years), and in the first year of primary school (i.e., they had already acquired literacy skills; mean age = 6.32 years). They used four-letter strings in which, for “different” trials, they had transposed-letter and replaced letter pairs. Duñabeitia et al. (2015) only found a transposed-letter effect when the children were in first grade (42.9% vs. 30.6% of errors, for transposed vs. replaced-letter pairs) and concluded that “position uncertainty emerges as a consequence of literacy training” (p. 549). However, the null effects obtained with the children in their pre-school years faced interpretative difficulties because these children did not adequately perform the task (d' was close to zero; see Perea et al., 2016). Indeed, Perea et al. (2016) found robust transposed-letter effects in pre-literate children with a simplified version of the same-different task. Nevertheless, they did not run a re-test when these children learned to read—note that first-graders could have

been at ceiling with this simplified task, thus making the comparison uninterpretable.

To test the emergence of location-invariant processing, while avoiding the interpretive difficulties of comparing performance of pre-school vs. school children, we conducted a same-different matching task with adult participants using transposed-letter vs. replaced-letter pairs as “different” trials before and after learning to read in a new script (Experiment 1). To control for mere visual familiarity, we included an untrained script that was also presented during training. Results showed transposed-letter effects of similar size for the trained and untrained scripts in both the pre-and post-training phases. Furthermore, Bayesian analyses offered substantial evidence in favor of a null interaction between training, phase, and script. Thus, learning to read and write in a new script does not lead to the rapid emergence of location-invariant processing.

Importantly, we did find some training benefit in the trained script for “same” responses in both response times and accuracy, which suggests the emergence of an early and basic visual specialization for letter strings. However, the newly acquired expertise in the new script was not sufficient to induce location-invariant processing. Indeed, the transposed-letter effect was greater for the Roman script than for the trained script (i.e., 28.3% vs. 19.2% of errors, respectively). This is the typical pattern when comparing strings of letters vs. strings of other visual objects (e.g., symbols, unknown letters). This pattern can be parsimoniously explained in terms of an orthographically-specific location-invariant component in the Roman script over and above the location uncertainty common to all visual objects (see Massol et al., 2013).

Our findings can also shed some light on the early developmental trajectory of the letter-specific position coding when learning to read. The absence of the emergence of location-invariant processing in the very early stages of learning to read can be accommodated by the dual-route model of orthographic development proposed by Grainger and Ziegler (2011; see also Grainger et al., 2012; Ziegler et al., 2014). This model assumes that, in the first stages of reading acquisition, the processing of letters in a word is serial, thereby

letter position coding is very strict (i.e., fine-grained orthographic coding). It is only when readers have more extensive reading experience that a more parallel processing of letters is developed, thus speeding the mapping of letters onto orthographic representations and producing greater transposed-letter effects (i.e., coarse-grained orthographic coding; see Grainger et al., 2012). In the context of the current experiment, participants acquired some basic orthographic skills, as revealed by better performance for “same” responses in the post-training phase. However, this expertise did not suffice for a coarse-grained processing to emerge. Indeed, in a lexical decision experiment that compared the error rates to pseudohomophone and orthographic controls, Grainger et al. (2012) found effects greater than 30% in Grade 1 and Grade 2 children—these effects were smaller with older children (i.e., the effects were 20% in Grade 3, 21% in Grade 4 and 16% in Grade 5). That is, beginning readers use phonological recoding (i.e., a fine-grained orthographic coding) rather than the coarse-grained coding responsible for location-invariant processing. Thus, the greater transposed-letter effect in the Roman (overlearned) script than in the newly learned script obtained in the current experiments suggests that the emergence of location-invariant processing requires a more complete establishment of a written orthographic code (see Grainger et al., 2012; Grainger & Ziegler, 2011; Ziegler et al., 2014).⁹

The emergence of location-specific processing

The second research question was whether location-specific processing emerges rapidly for the newly learned script using a target-in-string identification task. The dissociation in the accuracy serial position functions of letters (W-shape

⁹ An alternative account of orthographic development is Castles et al.’s (2007) lexical tuning model. The model assumes that acquiring more and more words in the lexicon involves an increasingly dense neighborhood of orthographically similar words. To efficiently identify these words, the orthographic representations become increasingly fine-tuned—this includes more precise positional representation of the visual input. Our experiments, however, were not designed to test the development of the orthographic lexicon (i.e., participants were trained to read and write pseudowords).

function) and symbols (Δ -shape function) is this task is assumed to be to an adaptation of the mechanisms of visual object processing to cope with visual word processing (i.e., location-specific letter detectors creation). Importantly, Dandurand et al. (2010; Grainger et al., 2016) hypothesized that the conversion of the mechanisms of simple visual object processing into location-specific detectors occurs with reading acquisition.

Results in the target-in-string identification task (Experiment 2) showed a clear advantage of the middle position for both the trained and control scripts in all scenarios. More critically, we found no signs of an interaction in accuracy between training, phase, and position, as shown in the Bayesian analyses. We also found a small advantage of the initial letter position over the other letter positions (see Figure 6; see also Appendix A), but this difference was not modulated by training (i.e., the difference was approximately constant in the pre- and post-training phases and in the trained and untrained scripts). Finally, the overall accuracy in the post-training phase was greater than in the pre-training phase for both, the learned and the control script, but this occurred essentially for the middle, fixated, letter. This latter finding can be parsimoniously explained in terms of better performance due to increased visual familiarity rather than on location-specific processing.

To our knowledge, unlike for location-invariant processing, no study has directly examined the emergence of location-specific processing—neither with children nor with adults. Nevertheless, for comparison purposes, it may be relevant to briefly discuss the studies that examined the developmental trajectory of the first-letter advantage in the very early stages of learning to read. Grainger et al. (2016) showed a small increase in the initial-letter advantage across school grades (from 1st to 4th) (see also Schubert et al., 2017, for a similar pattern of results). Importantly, the accuracy in the first and second position for 1st-3rd grade children were very similar, around 55% and 60% in the Grainger et al.'s (2016) experiment. The lack of a sizeable first position advantage in the initial grades in developing readers is in consonance with the results of Experiment 2. Taken together, these findings suggest that location-specific processing does not

emerge in the first stages of learning to read; instead, there appears to be a long route for this mechanism to emerge and develop.

Nevertheless, the lack of the sizeable first-letter advantage in the (overlearned) Roman script suggests that some caution when interpreting the findings of Experiment 2. In the experimental setup, participants always received the blocks with the artificial scripts (i.e., the main blocks for the purposes of the experiment) before the final block with the Roman script and, furthermore, instructions stressed that they should fixate at the center position at the beginning of each trial. Thus, a parsimonious explanation of the lack of a substantial first-letter advantage in the Roman script is that, to cope with the highly demanding blocks with the artificial scripts, participants' attention was focused on the middle position and this strategy was dragged into the Roman block. Thus, one might argue that the settings of Roman block were not optimal to capture a W-serial position function in the Roman block. Future research should examine to what degree the accuracy function in target-in-string identification tasks is modulated by task context and instructions (e.g., see Winskel et al., 2014, for evidence of different accuracy functions depending on the nature of the writing system).

On the emergence of orthographic processing when learning to read a new script

Recent research with adult readers has shown that orthographic processes can emerge rapidly after learning a script during a relatively short amount of time (e.g., Chetail, 2017; Taylor et al., 2011; Lally et al., 2020). For instance, in the Taylor et al.'s (2011) experiments, adults learned 36 words in an artificial script during 30-45 minutes. In the post-training phase, participants had to discriminate between trained and untrained items (i.e., an analogue to lexical decision). Results showed that participants could successfully discriminate trained from untrained items and, more important, the response times to the trained items were sensitive to vowel frequency. In a subsequent generalization phase, participants were asked to read aloud a series of new (untrained) items. Results showed an effect of both vowel frequency and consistency. All and all, the Taylor

et al. (2011) experiments suggest that participants can quickly and efficiently extract sub-word spelling-sound regularities in a new script (see Chetail, 2017, for a similar pattern of results regarding letter and bigram frequency).

More recently, Lally et al. (2020) conducted an experiment in which participants learned 24 five-letter pseudowords either in a sparse or in a dense artificial orthography, using a between-subject design, during a four-day training. (The 24 pseudowords in the dense orthography included 12 anagram pairs, whereas none of the 24 pseudowords sparse included anagrams.) When tested in an old-new recognition task (i.e., they had to discriminate between trained and untrained items), participants made fewer false positives for untrained items created by transposing two letters in the dense orthography than in the sparse orthography. Therefore, the findings reported by Lally et al. (2020) are a demonstration that the properties of the writing systems may modulate how letter order is encoded in a newly learned script.

In the current experiments, participants were able to read and write in the trained script with some fluency. Notably, in line with above-cited studies with artificial script training, we found some letter-specific processing as a consequence of learning to read: responses to “same” trials in the same-different task were faster and more accurate in the post-trained phase for the trained script, but not for the untrained script. As Krueger (1978) indicated, fast and accurate responses for “same” trials imply that participants require an exhaustive processing of the letter string. Thus, this pattern suggests that early literacy induced some specialization for letter strings that made the identification of same pairs more effortless. However, we found no signs reflecting the emergence of location-invariant and location-specific processing in the newly learned script.

In addition, we found an improvement in performance in the post-training phase for both the learned script and the visual control script that can be parsimoniously explained in terms of visual familiarity. Indeed, previous event-related potentials (ERP) experiments have shown that the N1 component (i.e., a component related to familiar written language processing) emerges right-lateralized in preschoolers at the start of reading training (e.g., Maurer et al.,

2006). Crucially, this early emergence of the N1 effect has been associated with letter knowledge and it also likely reflects visual familiarity with print (Maurer, Brem et al., 2005). The development of the characteristic left-lateralization of the component is assumed to occur via the automatization of orthographic-phonological mappings established during learning to read. (Maurer & McCandliss, 2007; McCandliss & Noble, 2003; Maurer et al., 2006; Posner & McCandliss, 2000; see also Maurer et al., 2010, for evidence with adults learning an artificial script; Brem et al., 2005, for the same pattern with visual training of symbols). Thus, the increase in performance in the post-training phase, coupled with the absence of differences in location-invariant and location-specific processing between the two phases, favors the idea that these effects are associated to visual expertise with the novel scripts (i.e., a reading-related perceptual expertise; see Maurer et al., 2010, for similar claims).

What we should also note is that, although both the same/different matching task and target-in-string identification task have been widely used to demonstrate orthographic effects (e.g., see Duñabeitia et al., 2012; Massol et al., 2013; Perea et al., 2016; Scaltriti & Balota, 2013; Tydgate & Grainger, 2009), they can be performed on the basis of visual representations of the stimuli—this is the reason why these tasks can be used in both pre-training and post-training phases. As a result, some effects obtained from newly learned scripts may reflect a mixture of increased visual familiarity to the new letters together with some incipient orthographic representations (i.e., a specific to letters visual expertise, see Maurer, Brandeis, & McCandliss, 2005). Instead, skilled readers, who have already automatized the orthographic-phonological mappings, would perform these tasks not only on the basis of visual familiarity, but also on the activation of the orthographic representations of the stimuli. In other words, the patterns observed in beginner readers may rely on visual familiarization with the learned scripts, whereas the effects of skilled readers may represent an interaction between early visual processing at the letter level and feedback from orthographic representations (see Marinus et al., 2018).

Taken together, our findings suggest that in order to boost the automatization of the orthographical-phonological mappings and a more parallel

coarse-coding processing, a much longer learning-to-read period may be required. We now discuss several options for further research. The first option would be to run a large-scale longitudinal experiment with pre-literate children—for the sake of the argument, we assume that the experimental tasks would allow a meaningful comparison across age (see Perea et al., 2016, for discussion). The experimental design would include three scripts: Roman letters, digits (i.e., a familiar visual object), and a control artificial script. This would allow us to examine not only the emergence of location-invariant and location-specific processing in a natural setting (i.e., children learning to read), but also how these orthographic markers vary as a consequence of literacy acquisition. Furthermore, this design would also allow examining the variations due to orthographic processing (i.e., specific to letters) vs. visual familiarity (numbers vs. artificial letters). A second option would be to train adult participants for a long period of time in an ecological setting—instead of the ecological limitations of learning an artificial script. The most realistic design would be to run the experiment with adults who are starting to learn a new language that uses an unfamiliar alphabetic script (e.g., Georgian, Armenian). In either scenario, it would be desirable to complement the behavioral tasks with the recording of brain activity during training, as this may help to disentangle the effects due to orthographical-phonological decoding from the effects due to visual training (e.g., see Maurer, Brem, et al., 2005, 2006; Pleisch et al., 2019, for evidence of print sensitivity in the N1 ERP amplitude; see Pleisch et al., 2019, for evidence of changes in the activation of crucial orthographic processing brain regions [ventral occipitotemporal cortex and left fusiform gyrus] in the first steps of reading acquisition).

To sum up, we conducted two experiments that examined whether two markers of orthographic processing (location-invariant and location-specific processing) arise rapidly after learning to read and write a new script. Notably, examining when these effects emerge is essential to help interpret the subsequent developmental trajectory of orthographic effects. While participants were able to read and write with some fluency in the new script and showed some rudimentary orthographic processing, we found no evidence favoring the

hypotheses that location-invariance and location-specific processing emerge quickly after learning to read. Instead, the emergence of these two markers of orthographic processing may take much more time, probably via the automatization of orthographic-phonological mappings.

Appendix A

Supplementary analyses with Bayesian linear mixed effects models

Experiment 1: location-invariant processing

For the sake of completeness, we also examined the latency and accuracy data of the confirmatory analyses using Bayesian linear mixed effects models using the *brms* package in R (Bürkner, 2016). An advantage of this procedure—via Stan—is that it allows us to fit the models using the maximal random effect structure (see Bates et al., 2015, for arguments in favor of maximal models).

Different trials

Trained script vs. Untrained script. The fitted model was:

```
Dependent_Variable ~ script * condition * phase + (1 + script *  
condition * phase | subject) + (1 + condition * phase | item)
```

More complex random effects terms resulted in model nonconvergence. Furthermore, these models offer the Bayesian 95% credible intervals for each parameter based on the posterior distributions. For the latency data, we employed the same response time transformation as in LME analyses (i.e., $-1,000/\text{RT}$; family = gaussian), whereas for the accuracy data, we used the Bernoulli distribution (family = bernoulli)—this is the parallel to family = binomial in GLME models. We employed 4 chains, each with 10,000 iterations after a warmup of 1000 iterations. The maximal random effect structure models converged successfully: the values of $\hat{R}$ were 1.00 for all parameters.

As can be seen from the estimates and 95% Credible intervals presented in Table A1, we found robust evidence of an effect of Condition (i.e., probe-target relationship) and Phase in both latency and accuracy analyses, thus corroborating the pre-registered analyses. Similarly, we did not find any evidence of a three-way interaction between Script, Condition, and Phase (see Figure A1, for the posterior distributions).

Table A1

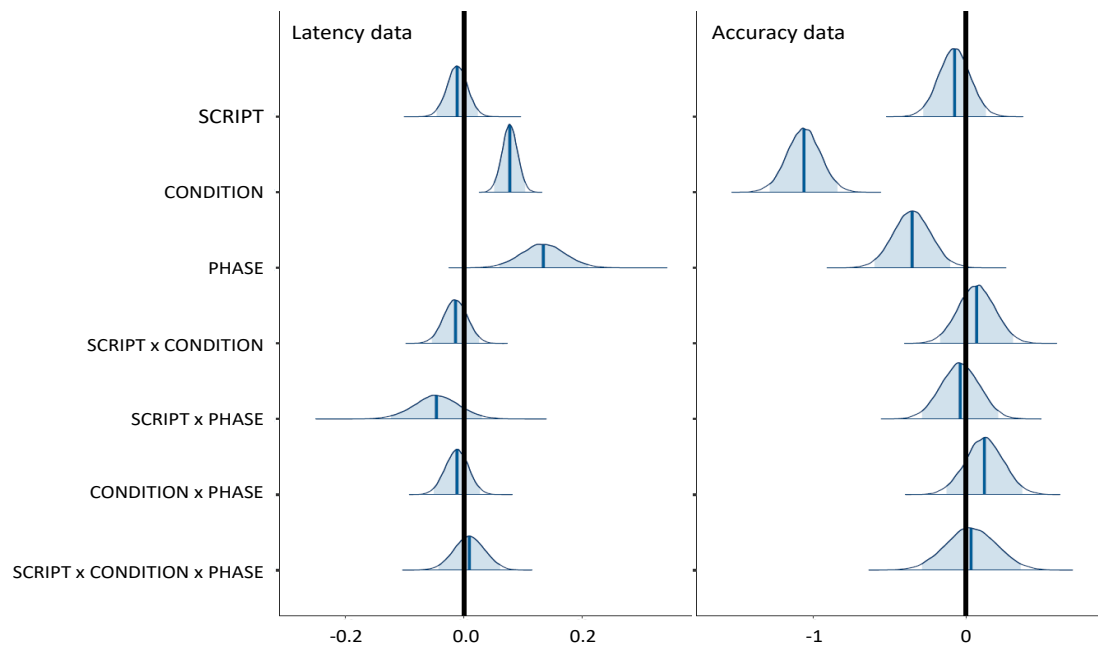
Parameter estimates in latency and accuracy supplemental analyses of Experiment 1 (“different” trials: trained vs. untrained script)

Latency Data	Estimate	SE	95% Credible Interval
Intercept	-1.84	0.06	[-1.95, -1.73]
Script	-0.01	0.02	[-0.05, 0.02]
Condition	0.08	0.01	[0.05, 0.10]
Phase	0.13	0.04	[0.06, 0.21]
Script x Condition	-0.01	0.02	[-0.05, 0.03]
Script x Phase	-0.05	0.04	[-0.12, 0.03]
Condition x Phase	-0.01	0.02	[-0.05, 0.03]
Script x Condition x Phase	0.01	0.03	[-0.04, 0.06]
<i>Accuracy Data</i>			
Intercept	1.56	0.15	[1.27, 1.87]
Script	-0.07	0.10	[-0.28, 0.13]
Condition	-1.06	0.11	[-1.28, -0.84]
Phase	-0.35	0.12	[-0.60, -0.10]
Script x Condition	0.07	0.12	[-0.17, 0.31]
Script x Phase	-0.04	0.13	[-0.29, 0.21]
Condition x Phase	0.12	0.13	[-0.13, 0.37]
Script x Condition x Phase	0.03	0.17	[-0.29, 0.36]

Note: Those effects with 95% credible intervals beyond 0 are in bold.

Figure A1

Posterior effects estimates from the Bayesian linear mixed models for “different” trials in Experiment 1 (Trained vs. Untrained script) (left panel: latency analysis; right panel: accuracy analysis)



Note: The thick black line corresponds to an effect of zero, the dark grey line corresponds to the estimates and the shaded area corresponds to the 95% credible interval.

Roman script vs. trained script. The fixed factors were Script (Trained [post-training] vs. Roman) and Condition (transposed, replaced). We followed the same procedure as above, with 5,000 iterations ($\hat{R} = 1.00$ in all cases). Table A2 showed the estimates and 95% credible intervals from the latency and accuracy models and Figure A2 showed the posterior distributions of the parameters. Together with a substantial transposed-letter effect in the latency data, these analyses confirmed the interaction between Script and Condition in accuracy data.

Table A2

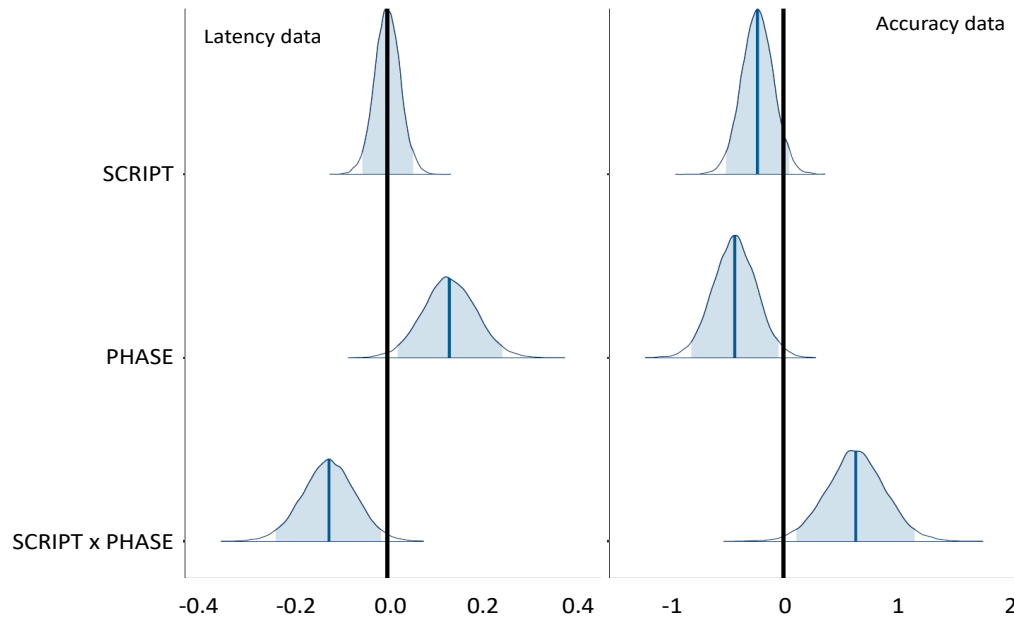
Parameter estimates in latency and accuracy supplemental exploratory analyses of Experiment 1 (“different” trials: Roman vs. trained script)

<i>Latency Data</i>	Estimate	SE	95% Credible Interval
Intercept	-1.83	0.04	[-1.91, -1.74]
Condition	0.07	0.02	[0.04, 0.11]
Script	-0.02	0.05	[-0.13, 0.09]
Condition x Script	0.00	0.02	[-0.04, 0.05]
<i>Accuracy Data</i>			
Intercept	1.86	0.17	[1.55, 2.20]
Condition	-1.66	0.14	[-1.92, -1.39]
Script	-0.26	0.21	[-0.69, 0.16]
Condition x Script	0.56	0.17	[0.23, 0.90]

Note: Those effects with 95% credible intervals beyond 0 are in bold.

Figure A2

Posterior effects estimates from the Bayesian linear mixed models in “different” trials in Experiment 1 (Roman vs Trained script) (left panel: latency analysis; right panel: accuracy analysis).



Note: The thick black line corresponds to an effect of zero, the dark grey line corresponds to the estimate and the shaded area corresponds to the 95% credible interval.

Same trials

Trained script vs. Untrained script. The general procedure was the same as above (i.e., the maximal random effect structure model with 5,000 iterations; $\hat{R} = 1.00$ in all cases). As shown in Table A3 (estimates and 95% credible intervals) and Figure A3 (posterior distributions of the parameters), these analyses corroborated the interaction effect for between Phase and Script in the latency and accuracy data.

Table A3

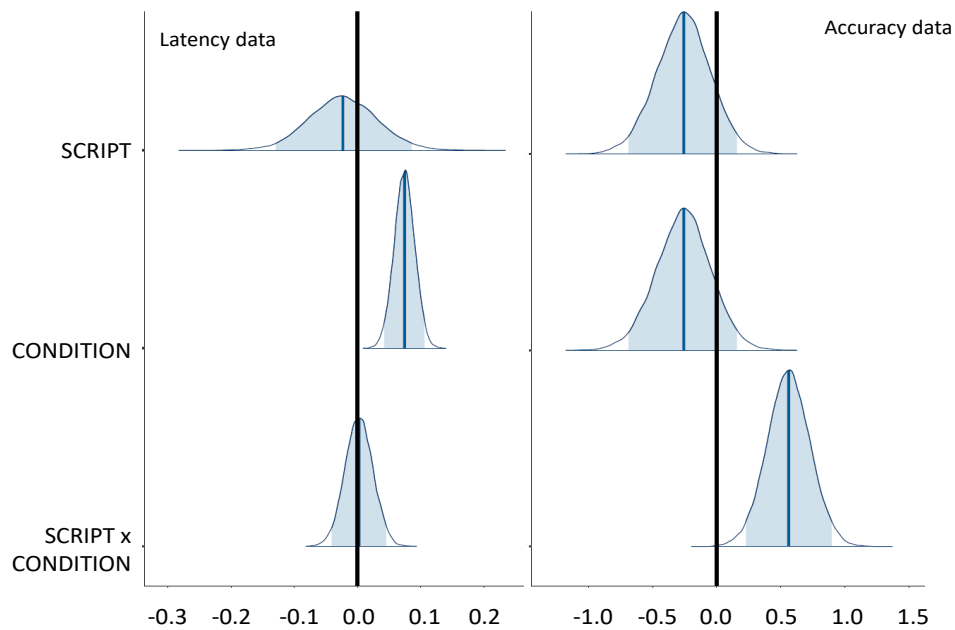
Parameter estimates in latency and accuracy supplemental analyses of Experiment 1 (“same” trials: trained vs. untrained script)

<i>Latency Data</i>	Estimate	SE	95% Credible Interval
Intercept	-1.99	0.08	[-2.14, -1.84]
Script	0.00	0.03	[-0.05, 0.05]
Phase	0.13	0.06	[0.02, 0.24]
Script x Phase	-0.13	0.06	[-0.24, -0.01]
<i>Accuracy Data</i>			
Intercept	2.70	0.21	[2.31, 3.12]
Script	-0.23	0.14	[-0.51, 0.05]
Phase	-0.43	0.19	[-0.81, -0.05]
Script x Phase	0.64	0.26	[0.11, 1.15]

Note: Those effects with 95% credible intervals beyond 0 are in bold.

Figure A3

Posterior effects estimates from the Bayesian linear mixed models in “same” trials of Experiment 1.



Note: The thick black line corresponds to an effect of zero, the dark grey line corresponds to the estimates and the shaded area corresponds to the 95% credible interval.

Experiment 2: location-specific processing

Trained vs. Untrained script. As in Experiment 1, we examined the data with Bayesian linear mixed effects models using the *brms* package (Bürkner, 2016) in R. The three fixed factors were the same as in the pre-registered analyses. The initial letter was set as the reference level for the factor Position. We fitted the maximal random effect structure model

```
accuracy ~ script* position*phase + (1 + script*position*phase|subject)
+ (1 + position*phase|item))
```

using 4 chains, each with 5,000 iterations after a warmup of 1,000 iterations. The priors were the same as in Experiment 1. The model converged successfully ($\hat{R} = 1.00$ for all parameters).

As can be seen in Table A4, accuracy in the initial letter position was substantially lower than in the middle, fixated position. The accuracy advantage of the middle position increased in the post-trained phase, as deduced from the interaction with Phase. In addition, the parameter estimates showed some advantage of the initial position over the other letter positions (see Figure A4, for the posterior effects). Therefore, these analyses corroborate the findings obtained in the ANOVAs indicated in the pre-registered analyses.

Table A4

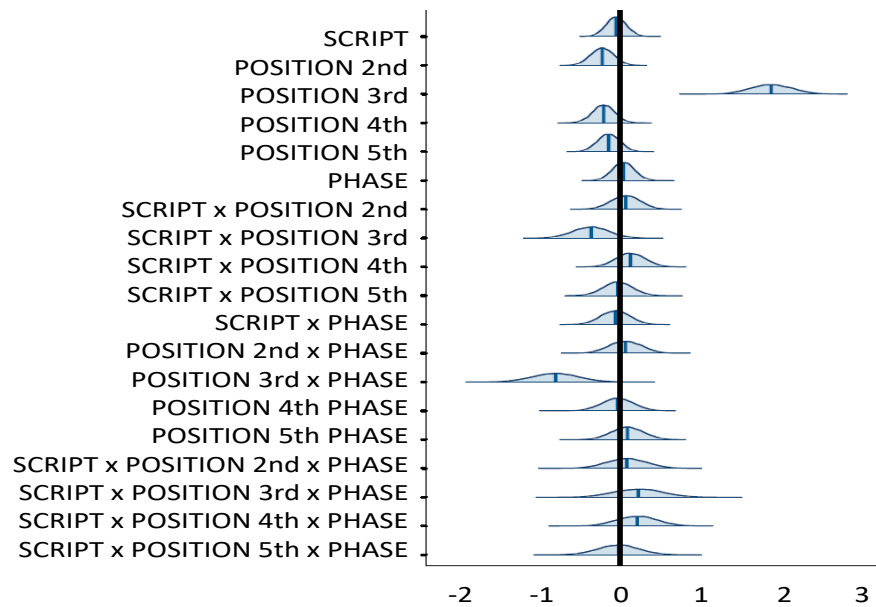
Parameter estimates in the accuracy supplemental analyses of Experiment 2 (trained vs. untrained script)

		Estimate	SE	95% Credible Interval
Intercept		0.16	0.09	[-0.03, 0.34]
Script		-0.05	0.13	[-0.29, 0.20]
Position	2 nd	-0.22	0.14	[-0.50, 0.06]
	3rd	1.89	0.26	[1.38, 2.41]
	4 th	-0.20	0.14	[-0.48, 0.07]
	5 th	-0.14	0.14	[-0.41, 0.13]
Phase		0.04	0.14	[-0.22, 0.31]
Script x	2 nd position	0.07	0.18	[-0.29, 0.42]
	3 rd position	-0.36	0.22	[-0.80, 0.08]
	4 th position	0.13	0.18	[-0.22, 0.49]
	5 th position	-0.03	0.18	[-0.39, 0.33]
Script x Phase		-0.06	0.18	[-0.41, 0.29]
Phase x	2 nd position	0.07	0.20	[-0.33, 0.46]
	3rd position	-0.80	0.29	[-1.38, -0.22]
	4 th position	-0.03	0.20	[-0.42, 0.37]
	5 th position	0.09	0.20	[-0.29, 0.49]
Script x Phase x	2 nd position	0.08	0.25	[-0.41, 0.58]
	3 rd position	0.23	0.31	[-0.38, 0.83]
	4 th position	0.21	0.25	[-0.28, 0.71]
	5 th position	-0.02	0.26	[-0.52, 0.48]

Note: Those effects with 95% credible intervals beyond 0 are in bold.

Figure A4

Posterior effects estimates from the Bayesian linear mixed models in Experiment 2



Note: The thick black line corresponds to an effect of zero, the dark grey line corresponds to the estimate and the shaded area corresponds to the 95% credible interval.

Roman script vs. learned script. The procedure was the same as above, except that the two fixed factors were Script (Roman vs. Trained) and Position—the reference level for the factor Position was also the first letter. The maximal model converged successfully ($\hat{R} = 1.00$ for all parameters). As can be seen from the 95% credible intervals (see Table A5; see also Figure A5, for the posterior distributions), we found a substantial advantage of the third letter position. In addition, there was a numerical advantage of the first letter position relative to the second letter position—this effect did not interact with script.

Table A5

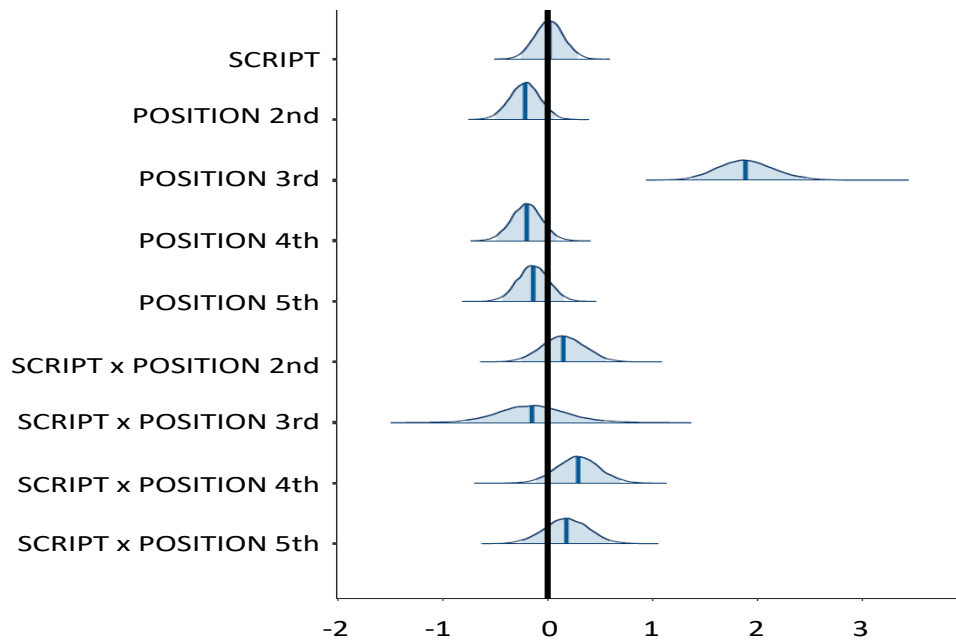
Parameter estimates in the accuracy supplemental analyses of Experiment 2 (Roman script vs. learned script)

		Estimate	SE	95% Credible Interval
Intercept		0.15	0.10	[-0.04, 0.35]
Script		0.01	0.14	[-0.25, 0.29]
Position	2 nd	-0.22	0.15	[-0.51, 0.06]
	3rd	1.90	0.28	[1.38, 2.47]
	4 th	-0.20	0.14	[-0.49, 0.08]
	5 th	-0.14	0.15	[-0.43, 0.15]
Script x	2 nd position	0.15	0.21	[-0.25, 0.55]
	3 rd position	-0.14	0.33	[-0.77, 0.53]
	4 th position	0.29	0.20	[-0.11, 0.69]
	5 th position	0.18	0.21	[-0.24, 0.59]

Note: Those effects with 95% credible intervals beyond 0 are in bold.

Figure A5

Posterior effects estimates from the Bayesian linear mixed models in Experiment 2 (Roman vs Trained script).



Note: The thick black line corresponds to an effect of zero, the dark grey line corresponds to the estimate and the shaded area corresponds to the 95% credible interval.

Appendix B

Figures B1 and B2 provide a visual representation of how training improved participants' performance along the learning days (from day 2 to day 6—note that on day 1 participants only had to listen to and write down the phoneme-grapheme correspondences).

Figure B1

Participants performance on the reading aloud and writing tasks along the training days (from Day 2 to Day 6).

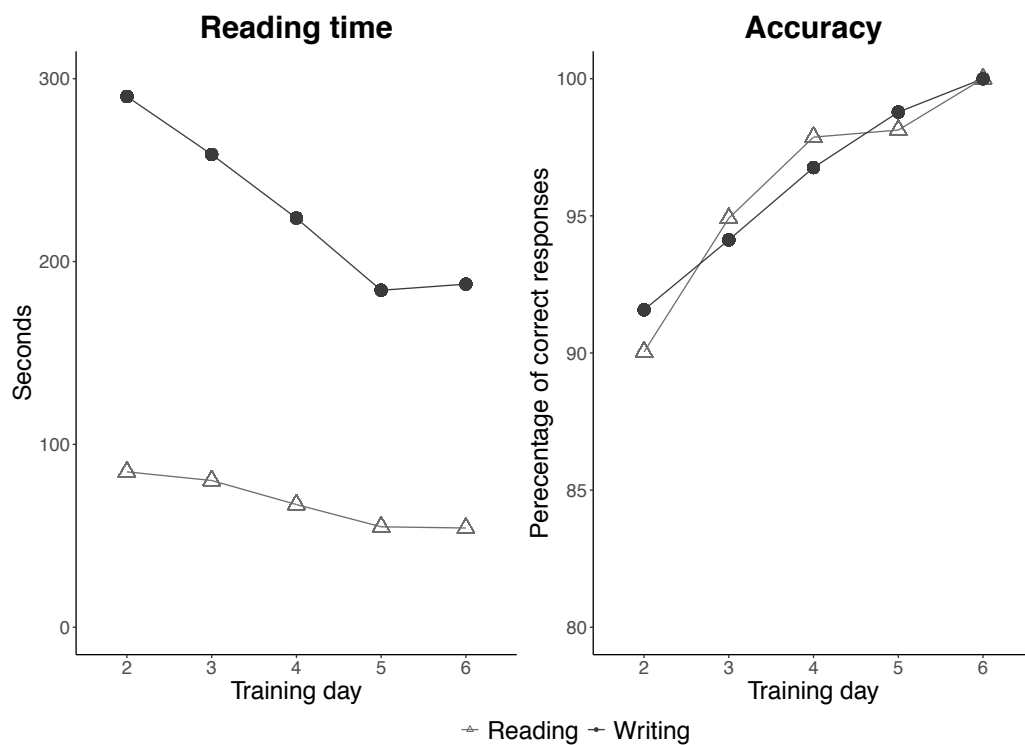
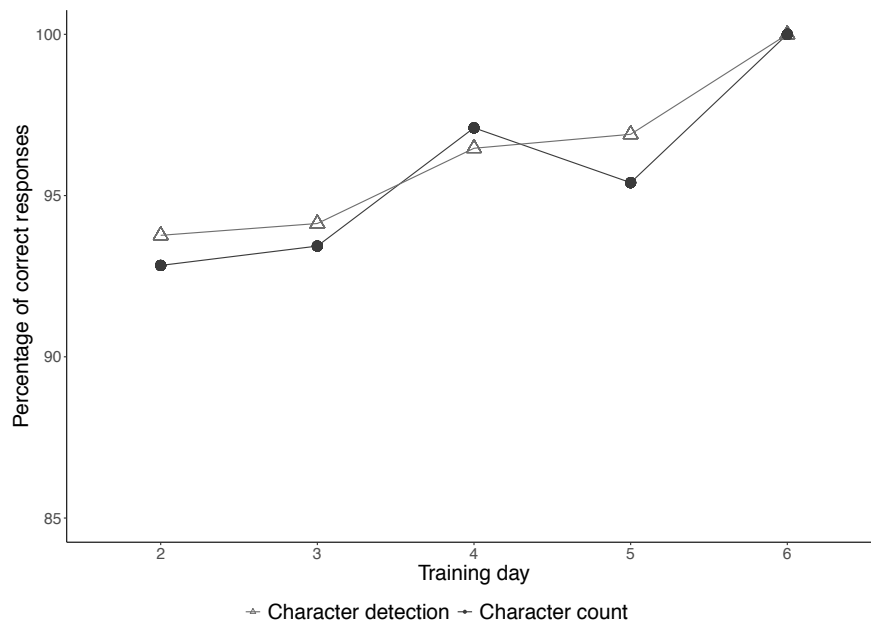


Figure B2

Participants performance on the visual familiarization tasks along the training days (from Day 2 to Day 6)



A letter is a letter and its co-occurrences: Cracking the emergence of location-invariance processing

Abstract

A central assumption of orthographic-based accounts of letter position encoding during word recognition is that, on top of position uncertainty, letter strings show location invariance (i.e., a key marker of orthographic processing; Grainger, 2018). One key finding supporting this view is that strings of letters are more vulnerable to transposition effects than strings of digits, symbols, or pseudoletters (e.g., the pair `CRFS-CFRS` is more confusable than `5732-5372`). In the present experiment, we tested whether location invariance emerges rapidly after the exposure to orthographic regularities—bigrams—in a novel script. To that end, we designed a study with two phases. In Phase 1, following Chetail (2017; Experiment 1b), individuals were first exposed to a flow of artificial words for a few minutes, with four bigrams occurring frequently. Afterward, participants judged the strings with trained bigrams as more wordlike (i.e., readers quickly picked up subtle new orthographic regularities) than the strings with untrained bigrams, replicating Chetail (2017). In Phase 2, participants performed a same-different matching task in which they had to decide whether pairs of 5-letter strings were the same or not. The critical comparison was between pairs with a transposition of letters in a frequent (trained) versus infrequent (untrained) bigram. Results showed that participants were more prone to make errors with frequent bigrams than infrequent bigrams with a letter transposition. These findings reveal that location invariance emerges rapidly after continuous exposure to orthographic regularities.

Key words: orthographic processing; orthographic regularities; letter position coding; artificial script

Introduction

If we consider reading using architectural terms, its building blocks are words, which, in turn, in alphabetical languages, are made up of letters. There is broad consensus that the identification of a written word is mediated by a process in which, from sensory input, the word recognition system encodes the abstract identities of letters in a specific order, thus allowing us to discriminate `hiss` from `kiss` and `dog` from `god`. In this context, orthographic processing constitutes the bridge between the sensory input and the word level (i.e., the encoding of letter identities and positions; see Grainger, 2018).

This paper focuses on the emergence of a fundamental marker of orthographic processing, location invariance, which comprises the mechanisms responsible for encoding the “relative positions of a set of object identities” (Grainger, 2018, p. 345). Previous research has shown that the way the human brain encodes the positions of letters in alphabetic languages is fairly flexible. For instance, `mohter` is easily confused with `mother`. Similarly, when participants are asked whether two successive strings of letters are the same or different, they respond “different” more slowly (and with less accuracy) to transposed-letter pairs (e.g., `CFLZ-CLFZ`) than to control, replacement-letter pairs (e.g., `CFLZ-CDVZ`), revealing a letter–transposition cost. One of the most common explanations for letter transposition effects is that they are due to positional uncertainty in assigning order between visual objects (e.g., letters). For example, the letters `H` and `T` in `mohter` would activate their own and neighboring positions. As a result, `mohter-mother` (or `CFLZ-CLFZ`) would be perceptually very similar (perceptual-based models; e.g., overlap model, Gomez et al., 2008; spatial coding model, Davis, 2010; Bayesian reader, Norris & Kinoshita, 2012). Of note, these accounts would also predict sizable transposition effects for other types of visual objects (e.g., `7586` would be perceptually similar to `7856`). A competing explanation is that letter transpositions effects are due to an orthographic mechanism that encodes location-invariant letter-in-word-order (orthographic-based models; e.g., open bigram model, Grainger & van Heuven, 2004; SERIOL model, Whitney, 2001). This second group of models assumes that reading

experience leads to the development of a layer of ordered pairs of letter co-occurrences (called “open bigrams”) between the letter and word levels. For instance, the word `mother` would be composed of the open bigrams `MO-MT-MH-ME-MR-OT-OH-OE-OR-TH-TE-TR-HE-HR-ER`, where `MO` would refer to “M to the left of O”. If two contiguous letters from `mother` are switched—as in `mohter`, 93% of the bigrams remain unchanged, thus explaining why `mohter` is easily confusable with `mother`—or `CLFZ` with `CFLZ`.

Although there has been some debate as to which family of models offers the best explanation of letter transposition effects (e.g., Norris & Kinoshita, 2012 vs. Whitney et al., 2011), both accounts may be necessary to fully understand letter position coding (see Grainger, 2018; Marcet et al., 2019). One key phenomenon supporting this “hybrid” view is that transposition effects are sizable for strings of numbers and symbols in same-different matching tasks (i.e., a finding that favors perceptual-based models) but their magnitude is smaller than that of strings of letters (i.e., an interaction that favors orthographic-based models) (Duñabeitia et al., 2012; see also Fernández-López et al., 2021; Ktori et al., 2019; Massol et al., 2013). Of note, the suggestion that both perceptual and orthographic-based elements affect letter transposition effects goes back to Estes (1975).

The existence of an intermediate level of letter groupings is consistent with experiments that showed that orthographic regularities (i.e., regularities learned incidentally throughout print exposure) in the form of bigrams modulate the assignment of letter position in letter strings and words. For instance, in a perceptual identification task with briefly presented stimuli, Rumelhart (1985) reported that participants tended to commit transposition errors for letter strings containing illegal bigrams, such as `praikc`—note that `KC` is not legal at the end of words in English. They would often report `praick` instead, which contains the legal bigram `CK`. Similarly, Frankish and Turner (2007) observed very high error rates for transposed-letter pseudowords like `sotrm` (base word: `storm`) in a lexical decision task, despite containing illegal bigrams—one might have thought that it would be easy to respond “nonword” based on this illegality (see also Frankish & Barnes, 2008; Perea & Carreiras, 2008, for converging evidence

using masked priming). Thus, following Rumelhart (1985), “the perception of a certain letter in a certain position depends on what we perceive in adjacent positions” (p. 725).

A key question arising from these findings is: what makes individuals particularly sensitive to orthographic regularities, such as bigram frequency, in a given language? Previous research has shown that readers quickly embrace sublexical regularities. To simplify the inherent complexity of reading, orthographic processing naturally relies on the regularities of the written system and capitalizes on statistical cues like bigram frequency (Cassar & Treiman, 1997; Chetail, 2017; Leloukiewicz et al., 2020; Mano & Kloos, 2018). Crucially, this ability emerges very rapidly throughout the exposure to print. A paradigmatic case is a study conducted by Pacton et al. (2001). They found that French readers were able, since very early in their development (i.e., six-year-old), to discriminate a word-like stimulus from a non-word-like stimulus based on their implicit learned knowledge on orthographic regularities. For instance, when comparing *ommer* vs. *ovver*, readers preferred *ommer* because *v* is never doubled in French. That is, the participants relied their decision on the frequency of the bigrams *mm* vs. *vv* (see also Doignon-Camus & Zagar, 2014; O’Brien, 2014). While the above findings are very informative, they suffer from an inherent drawback: the effects of orthographic regularities such as bigram frequency cannot be easily disentangled from other relevant factors that influence visual word recognition: pronounceability, familiarity, or orthographic neighborhood (Chetail, 2017). Keep in mind that the frequency of the bigrams in a given language cannot be manipulated but selected; thus, the design cannot be genuinely experimental.

To overcome the above limitation, a practical strategy is to use artificial scripts. This was the approach that Chetail (2017) followed with adult readers, testing what type of orthographic regularities emerges quickly with the exposition of print material in a novel script. Specifically, participants were first exposed to a flow of 5-letter words in an unfamiliar script—Phoenician alphabet—for approximately 9 minutes. There were four trained bigrams, so each word was made up of one of these bigrams, always in the same position (see Table 1).

Thereupon, in a wordlikeness task, participants were more likely to judge a new string as similar to the strings learned in the exposure phase if the string contained one of the trained bigrams in its position (i.e., a frequent bigram). In sum, a few minutes of exposition were enough to develop considerable sensitivity to bigram frequency. Chetail (2017) successfully replicated these findings in a second experiment in which participants learned 32 artificial words with a phonological form (i.e., the print-to-sound correspondences) before performing the task. Chetail (2017) concluded: “the statistical learning operating on the stream of artificial words made of a sequence of new shapes may be already oriented towards orthographic processing” (p. 118; see Vidal et al., 2021, for an alternative explanation). This study however did not test whether participants were prone to location-invariant encoding after learning the new orthography.

The main goal of the present experiment was to test whether location invariance, a critical marker of orthographic processing, emerges very early using a protocol parallel to Chetail’s (2017) Experiment 1b. In recent research, Fernández-López et al. (2021, Experiment 1) failed to obtain evidence of location-invariant processing in an experiment in which participants learned to fluently read and write in a new script across six training sessions. They conducted, both before and after the training, a same-different matching task where the different trials were created by transposing or replacing two adjacent letters. Results showed a similar pattern of letter transposition effects post-training and pre-training. The authors concluded that orthographic processing, at least in the form of location-invariant processing, does not emerge rapidly after learning a new script. However, they did not directly manipulate the sublexical properties of the stimuli. Here, we tested whether bigrams could be a key sublexical property helping the emergence of location invariance (see Grainger, 2018; Rumelhart, 1985).

Table 1.*Reproduction of materials used in Phase 1: replication of Chetail's (2017) protocol.*

N of trials		Examples of the stimuli		
Exposition phase				
80	ᠠᠳᠤᠫᠤᠯ	AB	_ _ _	
80	ᠯᠠᠭᠠᠨ	_ CE	_ _	
80	ᠭᠠᠨᠠᠨ	_ _ FG	_	
80	ᠠᠨᠠᠨᠠ	_ _ _ HI		
Wordlikeness phase		Critical item	Control item	
Familiarity condition	10	ᠠᠳᠤᠫᠤᠯ	ᠠᠳᠤᠫᠤᠯ	The critical item entailed one of the frequent bigrams in its corresponding position.
	10	ᠠᠳᠤᠫᠤᠯ	ᠠᠳᠤᠫᠤᠯ	
	10	ᠠᠳᠤᠫᠤᠯ	ᠠᠳᠤᠫᠤᠯ	The control item was composed by five infrequent letters
	10	ᠠᠳᠤᠫᠤᠯ	ᠠᠳᠤᠫᠤᠯ	
Position condition	10	ᠠᠳᠤᠫᠤᠯ	ᠠᠳᠤᠫᠤᠯ	The critical item entailed one of the frequent bigrams in its corresponding position.
	10	ᠠᠳᠤᠫᠤᠯ	ᠠᠳᠤᠫᠤᠯ	
	10	ᠠᠳᠤᠫᠤᠯ	ᠠᠳᠤᠫᠤᠯ	The control item entailed a frequent bigram in a non-corresponding position.
	10	ᠠᠳᠤᠫᠤᠯ	ᠠᠳᠤᠫᠤᠯ	
Letter frequency condition	10	ᠠᠳᠤᠫᠤᠯ	ᠠᠳᠤᠫᠤᠯ	The critical item entailed one of the frequent bigrams in its corresponding position.
	10	ᠠᠳᠤᠫᠤᠯ	ᠠᠳᠤᠫᠤᠯ	
	10	ᠠᠳᠤᠫᠤᠯ	ᠠᠳᠤᠫᠤᠯ	The control item entailed two frequent letters in their position, but not composing a frequent bigram.
	10	ᠠᠳᠤᠫᠤᠯ	ᠠᠳᠤᠫᠤᠯ	

Note. In Experiment 1b, Chetail (2017) used Phoenician alphabet—the database of artificial characters, BACS, was published in 2017. Bold is ours, to emphasize the frequent bigram.

The present study was composed of two phases. In Phase 1, we reproduced Chetail’s (2017; Experiment 1b) procedure—we replicated the main findings. The novelty of our work relies on the addition of Phase 2, including a same-different matching task with transposed-letter pairs. In this task, the probe was 5-letter string presented for 300 ms and was immediately followed, one line below, by a target that could be the same or different. The critical comparison was the following: pairs in which the target was created by transposing a trained bigram (e.g., **AB**; probe: **AB**UVX, target: BAUVX) versus pairs in which the target was created by transposing an untrained bigram (e.g., ZX; probe: ZXFGU, target: XZFGU)¹⁰. As is common in this paradigm, we also included replacement-letter pairs: We replaced a frequent bigram with two frequent letters (e.g., probe: **AB**UVX, target: **EC**UVX) or replaced an infrequent bigram with two infrequent letters (e.g., probe: XVFGU, target: MNFGU)—this comparison explores whether the replacement of a frequent bigram makes the pair less perceptually similar than the replacement of an infrequent bigram.

Thus, the present experiment aims to examine whether location invariance emerges very early, after exposure to print, with the acquisition of orthographic regularities. Two outcomes are possible: If location invariance emerges rapidly from the representations created by the trained bigrams in the new script (e.g., **AB**), the sequence BA (BAUVX) could be confusable with **AB** (**AB**UVX) because of (1) perceptual uncertainty and (2) **AB** having a mental representation. Instead, ZX (ZXFGU) would produce some activation on XZ (XZFGU) based on perceptual uncertainty alone. Therefore, it would be more difficult to respond “different” (i.e., more errors, longer response times) for those pairs involving the transposition of a frequent bigram like **AB** (position uncertainty + bigram activation) than an infrequent bigram like XZ (position uncertainty). This outcome would provide the first demonstration of the rapid emergence of location invariance.

Alternatively, the trained bigrams in the exposure phase may have not yet formed relatively stable representations to allow location invariance. If so, pairs

¹⁰ Bold and underlining indicated, in the examples, the frequent bigram and the letter manipulation, respectively (see Table 1 for examples of the employed stimuli).

with a letter transposition in trained or untrained bigrams would produce the same results (i.e., transposition errors would be based on perceptual uncertainty alone). This would suggest that the emergence of location invariance requires more extensive reading experience.

Method

Participants

Thirty-six undergraduate students from the University of València participated in the experiment. This sample size, the same as Chetail (2017, Experiment 1b), allowed us to have 1,440 observations per condition for the critical comparison of transposition of trained versus untrained bigrams, which is in line with Brysbaert and Stevens' (2018) recommendations. We also calculated Bayes Factors (BFs) to obtain a measure of the evidence for or against the effect—of note, we found conclusive evidence in favor of a difference between the transposition of frequent versus infrequent bigrams in the accuracy data ($BF = 112.16$)¹¹. All participants were native speakers of Spanish with normal or corrected vision and no history of reading or hearing disorders. They signed an informed consent form before participating in the experiment, and the study was approved by the Experimental Research Ethics Committee of the University of València. Participants received a small monetary compensation.

Materials

We used 21 letters from the BACS font to devise the stimuli (BACS1 and BACS2 serif font; Vidal et al., 2017)—note that these characters were matched to the Roman script in complexity, number of strokes, junctions, and terminations (see Table 1 for an illustrative example of the stimuli).

¹¹ Bayes Factors (BFs) were computed with JZS priors using *bfrms* (Singmann & Gronau, 2021) in R (R Core Team, 2021). BFs above of 3, 10 and 100 represent moderate, strong, and conclusive evidence, respectively.

For the *exposure phase*, we created 320 items of five characters, each including a critical bigram. Eight characters were used to devise the four frequent bigrams (frequent bigrams: $\alpha\beta$, $\gamma\delta$, $\tau\lambda$, and $\eta\iota$; that are represented in the body of the manuscript as **AB** ____, _ **CD** __, __ **FG** __, ____ **HI**; the infrequent bigrams were formed by the letters ϕ , λ , $\leq$, $\neq$, $\cup$, $\perp$, Δ , Φ , $\S$, $\supset$, $\varnothing$, $\in$, $\vdash$). Each frequent bigram occurred in a specific position. In this manner, we created 80 items with the critical bigram in positions 1 and 2 (**AB**KOR; bold is ours to facilitate the identification of the frequent bigram); 80 in positions 2 and 3 (R**CE**TN); 80 in positions 3 and 4 (ZX**FG**O), and 80 in positions 4 and 5 (UFK**HI**). We also created sixteen 5-letter strings with Roman characters to act as fillers. Importantly, each stimulus contained non-repeated characters. We created four different lists to counterbalance the position of the critical bigrams.

For the *wordlikeness phase*, as in Chetail (2017), we created 120 pairs of new stimuli (40 items per condition). For the familiarity condition, the critical item entailed one of the frequent bigrams in its corresponding position (based on the exposure phase; 10 items per position). The control item was composed by five random characters of the pool infrequent letters (e.g., UFK**HI** vs. BNPAO). In the position condition, the critical items were paired with control items that included the same critical bigram, but in a different position than in the trained items (i.e., UFK**HI** vs. **HI**BNP). Finally, in the letter frequency condition, the critical items were paired with control items that entailed two frequent letters in their frequent position, but they did not compose a critical bigram (UFK**HI** vs. MO**EG**B) (see Table 1, for an illustration of the materials created for each condition). As in the exposure phase, we created 4 different lists to counterbalance the position of the critical bigrams. We also created six 5-character string pairs to act as practice trials.

For the *same-different matching task*, we created 320 five-character string pairs (probe and target) in BACS font. All character strings were composed of non-repeated letters. There were 160 same pairs and 160 different pairs. For the same pairs, 80 contained a frequent bigram in its standard position (**AB**UVX—**AB**UVX), and 80 were composed of infrequent letters (OVNKM—OVNKM). For the

different pairs, 80 were created by transposing two letters, and 80 were created by replacing two letters—all of them contained a frequent bigram in its standard position. The transposed-letter pairs were created by transposing two adjacent letters in the target, that could be frequent (ABUVX–BAUVX; 40 pairs of items) or infrequent (XZFGU–ZXFGU; 40 pairs of items). The replaced-letter pairs were created by replacing two adjacent letters in the target, which could be frequent (ABUVX–ECUVX; the replacement letters were other frequent letters not constituting a frequent bigram [E from CE and G from FG]; 40 pairs of items) or infrequent (XVFGU–MNFGU; 40 pairs of items). Each manipulation occurred in 4 different positions (1st-2nd, 2nd-3rd, 3rd-4th, and 4th-5th)—there were ten pairs of items per position. To counterbalance the position of the frequent bigrams, we created four lists following a Latin square. For the practice phase, we created eight additional 5-character string pairs.

Procedure

Each participant performed the tasks of familiarization, exposure, wordlikeness (phase 1) and same-different (phase 2). The tasks of phase 1 paralleled those employed Chetail (2017). Participants were tested either individually or in groups of two in a quiet room. DMDX software (Forster & Forster, 2003) was used to display the sequence of stimuli and to register the timing/accuracy of the responses. All stimuli were presented in a monospaced font (15 pt BACS for the artificial strings; 15 pt Courier New for the Roman letters) in black on a white background. Response times were measured from target onset until the participant's response. The whole session lasted about 30-40 min.

The *familiarization task* consisted of introducing the 21 new letters, presenting them one by one in a computer screen. Participants were told to hand-copy the characters on a sheet of paper, without a time deadline. Immediately after, all the new letters were presented and participants could look at them as long as necessary.

In the *exposure phase*, the 320 artificial character strings were presented one by one in the center of the screen. Display duration and inter-stimuli interval were 1,200 and 500 ms, respectively. Participants were asked to carefully look at

the stream of stimuli. To ensure that participants were actually focusing of the letter strings, 5% of trials were fillers composed of five letters in Roman script (e.g., MNRLT). The participants were asked to respond to fillers by pressing the space bar.

In each trial of the *wordlikeness task*, a pair of stimuli was presented (critical and control items) on the screen. The critical item was on the left part of the screen in 50% of the trials and on the right part in the other trials. Participants were asked to decide which stimulus was the more similar to those presented in the exposure phase by pressing the corresponding key in the keyboard. They were asked to decide as soon as possible, although there was no time boundary.

In the *same-different matching task*, participants were told that they would be presented with two strings of characters and that they would have to decide if they were the same or not by pressing the “yes” and “no” keys. Participants were instructed to make this decision as quickly and as accurately as possible. A fixation point (*) was displayed for 500 ms in the center of a computer screen on each trial. Next, the fixation point was replaced by a probe, which was presented for 300 ms and positioned 3 mm above the center of the screen. Then, the target item appeared one line 3 mm below the center of the screen. The target remained on the screen until the response or 2,000 ms had passed.

Results

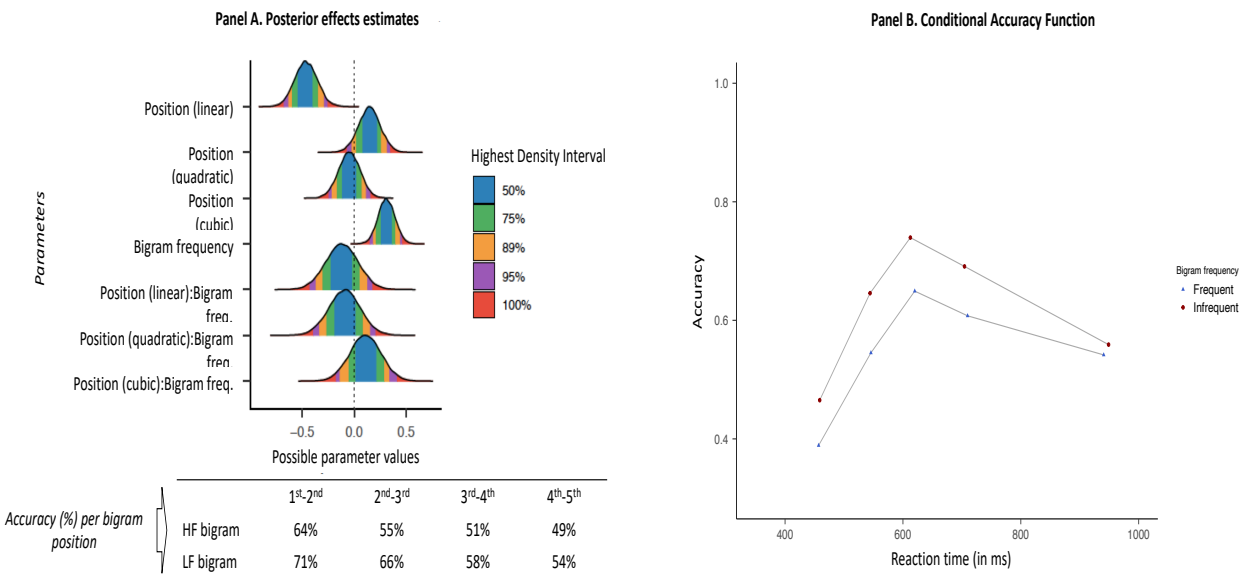
The analyses and results of Phase 1 are available in Appendix A—they essentially replicated Chetail’s (2017) findings: Participants were more likely to identify the foil containing previously learned bigrams as “wordlike” than that containing unfamiliar bigrams. For the inferential analyses of Phase 2 (same-different task), the dependent variables were the correct RT and accuracy. Very fast responses (< 250 ms: 9 responses) were omitted from the analyses of the correct RTs. Following our research goal, we tested the effect of the bigram frequency on transposed-letter processing—the analyses of “same” pairs and replacement-letter pairs are available in Appendix B. We fitted the data with

Bayesian linear mixed-effects models using *brms* (Bürkner, 2021) in R (R Core Team, 2021). The fixed effect was bigram frequency (frequent vs. infrequent) with the maximal random effect structure model for subjects and items. We used the Gaussian distribution to model the latency data (-1,000/RT) and the Bernoulli distribution to model the accuracy data (1=correct, 0=incorrect). For each model, we employed 10,000 iterations in each of the four chains (2,000 warmup + 8,000 sampling). The chains converged successfully (all $\hat{R}_s = 1.00$). The output indicates the estimate of each effect (the mean of the posterior distribution), together with its standard error and 95% credible interval (95%CrI). We inferred evidence of an effect when its 95%CrI did not include zero—this was complemented by its BF.

The models showed that accuracy was noticeably lower when the letter transposition involved a frequent than an infrequent bigram (54.7 vs. 62.0, respectively; $b = 0.33$, $SE = 0.09$, $95\%CrI [0.15, 0.50]$, $BF = 112.16$). The effect in the RTs had the same direction but was marginal (657 vs. 648 ms for the transpositions with frequent vs. infrequent bigrams), $b = -0.03$, $SE = 0.02$, $95\%CrI [-0.06, 0.01]$, $BF = 0.14$.

Figure 1.

Panel A: posterior effects estimates from the Bayesian linear mixed models. Panel B: conditional accuracy functions.



Note. In panel A, the dashed line corresponds to an effect of zero. In Panel B, the points represent the accuracy and average RT of the responses within equal sized bins (20% of responses per bin).

To scrutinize the effect of bigram frequency on transposed-letter pairs, we conducted two additional analyses. In the first analysis, we added position of the transposed letters (1st-2nd, 2nd-3rd, 3rd-4th, 4th-5th) as a polynomial factor in the design when analyzing accuracy. Results replicated the effect of frequency ($b = 0.31$, $SE = 0.08$, 95%CrI [0.15, 0.47]) and also revealed a linear component of position (accuracy decreased with position; $b = -0.46$, $SE = 0.11$, 95%CrI [-0.68, -0.25]) with no signs of an interaction (see Figure 1, panel A, for the descriptive statistics and posterior distributions). Second, we employed conditional accuracy functions (Figure 1, panel B), which describe how the response accuracy for a given condition varies across response speed. This allows us to examine whether the effect of bigram frequency occurred across the entire range of RTs or was limited to early or late responses. As shown in Figure 2, the conditional accuracy functions are approximately similar in the two conditions: regardless of speed response, responses were less accurate for the transposition of frequent bigrams than infrequent bigrams. Furthermore, both functions followed an inverted-U, where the fastest responses and, to a lesser degree, the slowest responses, were less precise.

Discussion

The present study examined whether a key marker of orthographic processing, location invariance, emerges quickly after the repeated exposition of orthographic regularities—bigrams—in a novel script. To that end, we designed an experiment that followed, in its first phase, the same procedure as Chetail (2017; Experiment 1b). For around 9 minutes, participants repeatedly received a series of 5-artificial-letter strings that contained a frequent bigram. Then, in a wordlikeness task, participants judged the strings with a trained bigram as more wordlike, thus showing that readers picked up subtle new orthographic regularities very rapidly, replicating Chetail (2017).

In the second—novel—phase of our experiment, participants performed a same-different matching task in which they had to decide whether pairs of 5-artificial-letter strings were the same or not. The critical test was whether the pairs involving the transposition of the letters of a frequent bigram (ABUVX–BAUVX) had a boost in confusability (i.e., location invariance in addition to position uncertainty) than those pairs involving the transposition of an infrequent bigram (XZFGU–ZXFGU) (i.e., position uncertainty). Responses to pairs with a letter transposition in a frequent bigram were less accurate than those with a letter transposition in an infrequent bigram (54.7 vs. 62.0%, respectively), thus showing an increase in confusability—the latency data were in the same direction¹². Critically, this is the first demonstration of the rapid emergence of location invariance in a novel script with an experimental design (see Rumelhart, 1985, for comparable evidence with legal vs. illegal bigrams in English).

Altogether, these findings favor the view that print exposure alone facilitates the development of orthographic regularities (see Chetail, 2017; Chetail & Sauval 2021). Critically, the repeated exposure to patterns of letter co-occurrences would facilitate the development of internal representations of letter clusters (i.e., frequent bigrams), inducing location invariance. As a result, when individuals are presented with BAUVX, the cognitive system would often confuse it with ABUVX because (1) BA has not occurred before, and AB has a mental representation, and (2) there is positional noise during order assignment. In contrast, only perceptual noise would affect order position in the strings that involved infrequent bigrams, such as XZFGU and ZXFGU. Thus, the combination of these mechanisms can readily explain why BAUVX is more confusable with

¹² Transposition effects in the same-different matching task are typically more robust for accuracy than for RTs (e.g., Duñabeitia et al., 2012; Fernández-López et al., 2021; Massol et al., 2013). This is not surprising when considering that the relatively high number of errors for transposed-letter pairs makes the latency data more variable and noisier than in paradigms with close-to-ceiling performance.

ABUVX than XZ**FGU** is with ZX**FGU**. A parallel rationalization would also apply to the confusability of praikc with praick (Rumelhart, 1985).¹³

The idea that readers internalize statistics on letter co-occurrences due to orthographic regularities aligns with the assumptions of orthographic-based models (Grainger & van Heuven, 2003; Whitney, 2001). However, we must also consider that the high number of errors for the pairs involving the transposition of an infrequent bigram (38%) strongly suggests that perceptual uncertainty is an inherent property of serial order encoding (Gomez et al., 2008; Norris & Kinoshita, 2012). Thus, the present findings add new evidence favoring the idea that both perceptual and orthographic-based processes underlie the effects of letter transposition (see Grainger, 2018). While this view is not new (see Estes, 1975), most of the evidence supporting orthographic-based models comes from a single phenomenon: the stronger transposition effects for letters than for digits and symbols in the same-different task.

Importantly, our findings also shed light on the early developmental trajectory of orthographic processing when learning to read. Firstly, statistical learning would facilitate acquiring orthographic regularities, such as frequently co-occurring letter combinations, that support the subsequent processing of higher-level linguistic entities. This idea fits well with the fact that preschoolers become rapidly sensitive to bigram frequency because of its functionality in learning to read (Mano & Kloos, 2018). Secondly, the extra confusability of transposed-letter pairs with frequent bigrams suggests that, in the first moments of learning to read, location-invariant processing is tuned to the processing of frequent letter chunks, helping the subsequent encoding of words. That is, the repeated exposure to **AB** clues us in that the bigram is probably present in many new to-be-learned words, so the location-invariant processing is tuned to encode **AB** for **AB** and **BA** with an optimization purpose. Hence, orthographic learning

¹³ Notably, although the orthographic regularities modulated the encoding of the order of letters, they did not have an effect in the encoding of letter identity: replacing a frequent or an infrequent bigram produced similar results.

would involve optimizing the mapping of letter-level information onto higher-level representations.

To sum up, the present findings demonstrated that orthographic regularities help the rapid emergence of location invariance when exposed to a new script. Importantly, this mechanism has an adaptive purpose: to help encode the to-be-learned words. In short, the ability to encode the properties of sublexical orthography may represent a unique ability within reading development, thus opening a window to further experiments examining sensitivity to orthographic regularities in early childhood.

Appendix A

Analyses of Phase 1

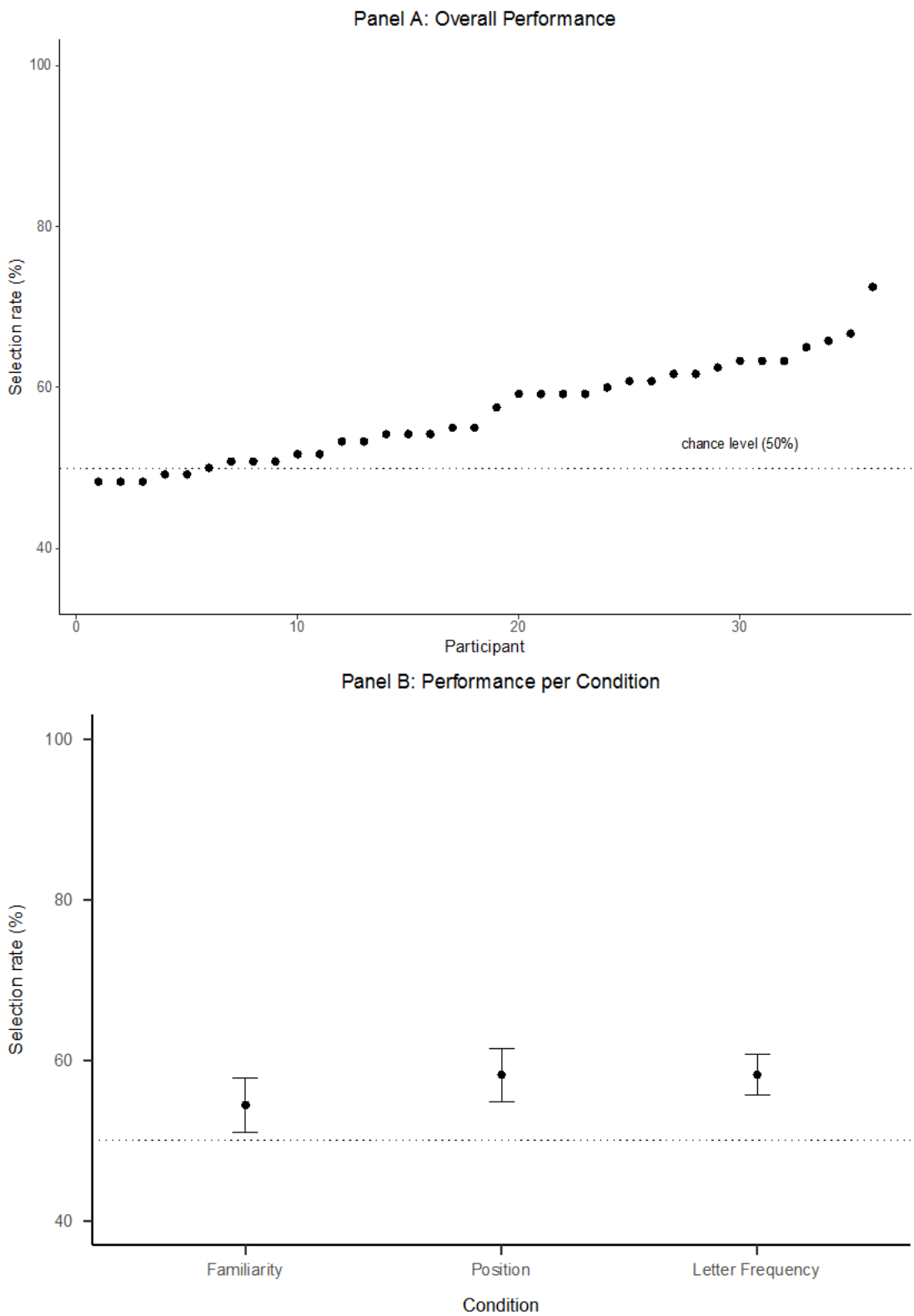
For the inferential analyses of Phase 1, we reproduced the analyses of Chetail (2017) and employed generalized linear mixed effects (GLME) models in R (R Core Team, 2021) using the *lme4* 1.1-27.1 package (see Bates et al., 2015) and the *lmerTest* package (Kuznetsova et al., 2017). The dependent variable was accuracy and we fitted GLME models with no fixed factor (comparison with chance level in the wordlikeness task: `glmer(ACCURACY~1+(1|PARTICIPANT))`). As a short teaser, we essentially replicated the findings.

Exposure: The detection rate of fillers was 99.74%.

Wordlikeness: Overall, participants performed above chance level ($M = 56.94\%$, see Figure A1, panel A), $b = 0.28$, $SE = 0.04$, $z = 6.71$, $p < 0.001$. The performance was significantly higher than chance level in the familiarity condition ($M = 54.44\%$), $b = 0.18$, $SE = 0.06$, $z = 2.63$, $p = 0.008$, in the position condition ($M = 58.19\%$), $b = 0.34$, $SE = 0.07$, $z = 4.83$, $p < 0.001$, and in the letter frequency condition ($M = 58.19\%$), $b = 0.33$, $SE = 0.05$, $z = 6.19$, $p < 0.001$ (Figure A1, panel B). Furthermore, performance was significantly higher for the initial bigram than for the final one ($b = 0.24$, $SE = 0.08$, $z = 2.83$, $p = 0.005$).

Figure A1.

Panel A: individual performance in the wordlikeness task. Panel B: selection rate of critical items in the familiarity, position and Letter Frequency conditions.



Appendix B

Results of Replaced-letter trials and Same trials in the Same Different task

Replaced-letter trials. Responses were only 3 ms faster for the frequent bigram replacements than for the infrequent bigram replacements (613 vs. 615 ms; $b = 0.00$, $SE = 0.02$, $95\%CrI [-0.03, 0.04]$). Moreover, there were no signs of differences in accuracy for the pairs with a replacement of frequent bigrams vs. infrequent bigrams (76.59 vs. 78.88, respectively; $b = 0.14$, $SE = 0.12$, $95\%CrI [-0.10, 0.37]$).

Same trials. There were virtually no differences in response times between pairs with frequent and infrequent bigrams (603 vs. 608 ms; $b = 0.01$, $SE = 0.01$, $95\%CrI [-0.02, 0.03]$). Regarding accuracy, accuracy was higher for the pairs containing a frequent bigram than for the pairs containing an infrequent bigram (89.72 vs. 87.55; $b = -0.33$, $SE = 0.14$, $95\%CrI [-0.62, -0.06]$).

Conclusions

One of the central landmarks of learning to read is the emergence of orthographic processing (i.e., the encoding of letter identity and letter order): it constitutes the necessary link between the low-level stages of visual processing and the higher-level processing of words. In this section, I presented three studies that aimed to fully understand the emergence of orthographic processing via the scrutiny of letter order and letter position coding.

In the first study, *Which factors modulate letter position coding in pre-literate children?* (Fernández-López et al., 2021), we aimed to explore the precursors of letter order coding. To that end, we examined the relationship between the ability of pre-literate children to accurately encode the order of letters (i.e., to differentiate between transposed-letter pairs and identity pairs: TZ–ZT vs TZ–TZ)—taken from the Perea et al. (2016) experiment—with the basic cognitive processes subtest of the BIL battery (Sellés et al., 2008). Results showed that the pre-literate children who best differentiated between transposed letter (TZ–ZT) and identity (TZ–TZ) pairs in a same-different task were those with higher scores on basic cognitive processes. Thus, basic cognitive skills shape the ability to encode serial order in letter strings.

In the second study, *Does orthographic processing emerge rapidly after learning a new script?* (Fernández-López et al., 2021) we aimed to track the emergence of early orthographic processes in the first steps of learning to read through the examination of location-invariant processing and location-specific processing. To that end, we employed a design with a pre-training phase and a post-training phase in which adults were trained in reading and writing nonsense words in an artificial script along six sessions. Participants were also visually familiarized with a different artificial script for control purposes. Regarding location-invariance, we run same-different matching tasks using transposed-letter vs. replaced-letter pairs as “different” trials before and after learning to read in a new script with the trained and control (visually familiarized) scripts. Results showed transposed-letter effects of similar size for the trained and untrained

scripts in both the pre-and post-training phases. Nevertheless, we did find that responses to same trials in the same-different task were faster and more accurate in the post-trained phase for the trained script, but not for the control script. Thus, learning to read and write in a new script did not lead to the rapid emergence of location-invariant processing. However, it seems that the learning-to-read training induced some rudimentary specialization for the letter strings that made easier the identification of same pairs. To test the emergence of location-invariant processing, we run target-in-string identification tasks with the trained and control scripts before and after learning to read. Results showed a clear advantage of the middle position for both scripts. Thus, learning to read and write in a new script did not lead to the rapid emergence of location-specific processing.

The emergence of orthographic processing probably needs much more time to emerge, or, alternatively it might require experience with specific orthographic structures. In the third study of the present section, *A letter is a letter and its co-occurrences: Cracking the emergence of location-invariance processing*, we tested whether location-invariant processing emerged rapidly after the exposure to orthographic regularities—bigrams—in a novel script. To that end, we designed an experiment with two phases. In Phase 1, we replicated the protocol of Chetail (2017; Experiment 1b), in which individuals were exposed to a flow of artificial words for a few minutes, with four bigrams occurring frequently. Then, in a wordlikeness task, they judged the strings with trained bigrams as more wordlike. Thus, readers quickly picked up subtle new orthographic regularities. In Phase 2, participants performed a same-different matching task. In contrast to the previous study (Experiment 1), where the critical test was transposed letter vs. replaced letter pairs, the critical comparison here was between pairs with a transposition of letters in a frequent (trained) vs. infrequent (untrained) bigram—note that our goal was to examine whether the transposition of the letters of a frequent bigram leads to a boost in confusability. Results showed that responses to pairs with a letter transposition in a frequent bigram were less accurate than those with a letter transposition in an infrequent bigram. This finding suggests that location-invariance emerges rapidly after exposure to orthographic regularities.

In summary, the three studies of this section aimed to further our understanding of the foundations of orthographic processing. Overall, it is worth highlighting the following premises: (i) Reading readiness emerges early in development with the adaptation of the mechanisms of visual object processing. For that reason, (ii) early detection of deficiencies in the visual analysis of the input is crucial to prevent later reading difficulties. Notably, we have added to the literature that (iii) orthographic processing does not emerge rapidly when learning to read, but, as a prelude, (iv) some non-specialized processes are fine-tuned to guide the later functional reading progress.

Thus, all things considered, the learning to read instruction should not be based solely on letter knowledge and phonological decoding. Teaching professionals should also contemplate that reading is built on a basic cognitive foundation that can be trained to smooth the learning-to-read. To name a few suggestions: training in a same-different task with letter stimuli can help with the encoding of letter order; and the mere visual exposure to the words would help create a rudimentary layer of frequent letter co-occurrences that facilitate the subsequent encoding of words.

The resilience of letter detectors

Words! Mere words! How terrible they were! One could not escape from them. They seemed to be able to give a plastic form to formless things. Was there anything so real as words?

Oscar Wilde, *The Picture of Dorian Gray*

Our brain is the outcome of millions of years of evolution in a world that, initially, did not contain letters—how can it adapt to the specific challenges posed by written word recognition? From the previous section, one can deduce that the brain is a piece of complex machinery subject to functional changes. Typically, when we think of modifying a machine, we imagine a person using a set of tools. The alphabet is one of these tools that literally transform the brain. Specifically, literacy changes the very subtle cognitive and brain mechanisms that connect visual forms with meaning.

As we know from the previous section, mere exposure to letters leads to some rudimentary functional changes that prepare our brain to read. Initially, reading is arduous and rigid, but because of the extensive experience across different conditions, skilled readers have developed some tolerance to noise and variations in the shape of the words' constituent letters (see Figure 1).

Figure 1.

Example of CAPTCHAs



Note. CAPTCHA is the acronym for Completely Automated Public Turing test to tell Computers and Humans Apart.

However, the empirical evidence on the resilience of the word recognition system to perceptual distortion is still scarce, probably because it entails many different elements (see Vergara-Martínez et al., 2021, for discussion). For instance, *handwritten text* is distorted in several dimensions, and hence any processing differences relative to the immaculately written format we typically encounter when reading might be due to multiple factors. In this section, I present three studies that aimed to examine the tolerance of the word recognition system to visual distortion in a single dimension: the rotations of letters.

An added value of studying letter rotations is that leading models of visual word recognition make explicit predictions about how letter rotations affect word reading. The SERIOL model of word recognition (Whitney, 2001, 2002) assumes that “input levels to letter units are reduced for rotated stimuli” (p. 119, Whitney, 2002), increasing gradually the time required to encode a letter string. Instead, the Local Combination Detectors model (Dehaene et al., 2005) offers an all-or-none account of how letter rotation affects letter and word processing. It postulates a boundary of 40°-45° beyond which letter detectors would not efficiently cope with rotations—there would only be a negligible cost for letter rotations below that boundary (see Cohen et al., 2008).

Nevertheless, this assumption of the LCD model has been recently questioned by several studies showing that readers can cope with rotations up to 180° from the initial access to orthographic representations. For instance, Perea et al. (2018) have found sizeable masked identity and transposed-letter priming effects for 90° rotated words, in line with the effects obtained with horizontally presented stimuli. This same pattern of results occurred for 180° rotated words in a study conducted by Yang and Lupker (2019). However, one can argue that as both primes and targets were rotated, participants could have developed conscious strategies, like tilting their heads, to cope with rotations. Thus, a suitable approach would be presenting the targets in the canonical orientation, preceded by rotated primes. This was the strategy followed by Benythe and Csibri (2021), who also found sizeable masked priming for rotated words with angles up to 90°.

Notably, at this point, we should consider that when a word is entirely rotated, readers can mentally rotate the stimulus as a whole, with one movement of mental rotation. By contrast, if we rotate the individual letters within words, individuals would need more than one rotation movement. Moreover, the rotation of the individual letters involves the disruption of trans-letter features. Thus, this manipulation allows us to examine the model predictions more stringently. For this reason, in the three studies of this section, we rotated the words on a letter-by-letter basis. Our main aim was to examine the resilience of letter detectors to different rotation angles and, in this manner, test the predictions of the SERIOL (Whitney, 2001, 2002) and Local Combination Detectors model (Dehaene et al., 2005).

In the first study, *How resilient is reading to letter rotations? A parafoveal preview investigation* (Fernández-López et al., 2021), we tested the resilience of letter detectors to rotations while reading. Specifically, we conducted a parafoveal preview experiment using the gaze-contingent boundary change paradigm. The parafoveal previews were identical or unrelated to the target word. Critically, each preview's individual letters were rotated 15°, 30°, 45°, or 60°. Apart from the parafoveal previews, all text, including target words, was presented with letters in their canonical upright orientation.

In the second study, *On the time course of the resilience of letter detectors to rotations: A masked priming ERP investigation*, we aimed to examine in detail the time course of processing words with rotated letters. To that end, we conducted a masked repetition priming lexical decision experiment while recording the participants' EEG measures. Targets were always presented in the standard horizontal format, and we rotated the individual letters of the identity/unrelated primes (0°, 45°, or 90°).

Finally, in the third study, *Letter rotations: Through the Magnifying Glass and What Evidence Found There*, we conducted three experiments to ascertain whether the impact of letter rotation was perceptual (visually) or linguistically mediated in an unprimed procedure. Experiments 1 and 2 employed four rotation angles (0°, 22.5°, 45°, 67.5°) in lexical decision and semantic categorization

tasks, respectively. In Experiment 3, we focused on four moderate rotation angles at or below 45° (22.5° , 30° , 37.5° , 45°).

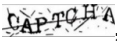
How resilient is reading to letter rotations? A parafoveal preview investigation

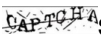
Abstract

Skilled readers have developed a certain amount of tolerance to variations in the visual form of words (e.g., CAPTCHAs, handwritten text, etc.). To examine how visual distortion affects the mapping from the visual input onto abstract word representations during normal reading, we focused on a single type of distortion: letter rotation. Importantly, two leading neurally-inspired models of word recognition (SERIOL model, Whitney, 2001; LCD model, Dehaene et al., 2005) make distinct predictions: whereas the SERIOL-model postulates that the cost of letter rotation increases gradually, the LCD-model assumes an all-or-none boundary at around 40°-45° rotation angle. To examine these predictions in a normal reading scenario, we conducted a parafoveal preview experiment using the gaze-contingent boundary change paradigm. The parafoveal previews were identical or unrelated to the target word. Critically, each preview's individual letters were rotated 15°, 30°, 45°, or 60°. Apart from the parafoveal previews, all text, including target words, was presented with letters in their canonical upright orientation. Results showed that the advantage of the identity preview condition in eye fixation times on the target word decreased progressively by rotation angle. This pattern of results favors the view that the cost of letter rotation during normal reading increases gradually as a function of the angle of rotation.

Key words: reading, parafoveal processing, eye-movements, letter rotation

Introduction

Efficient sentence reading requires the rapid abstraction of each word's visual form onto multiple levels of information (orthography, phonology, semantics, and syntax) while the eyes move along the text (see Clifton et al., 2016; Rayner, 2009, for reviews). Because of their extensive experience across different reading conditions (e.g., numerous *fonts*; uppercase TEXT, handwritten *text*, *cursive*, , etc.), skilled readers have developed some tolerance to noise and variations in shape (i.e., visual characteristics) of the words' constituent letters (see Grainger & Dufau, 2012).

While often overlooked in reading research, the study of the influence of visual factors in reading and how skilled readers tolerate various kinds of shape distortion is essential for a complete understanding of the reading process (Slattery 2016; see also Tinker, 1958; Tinker & Paterson, 1931, for reviews of early research). One of the reasons of the shortage of studies on this issue is that it is difficult to separate the effect of all the relevant variables underlying visual distortion. For instance, *handwritten text*, and s are distorted in several dimensions, and hence any processing differences relative to the pristine written format we typically encounter when reading might be due to multiple factors.

To overcome this interpretative issue, in the present paper, we focus on a single type of visual distortion during sentence reading: the rotation of individual letters in words. An added value of studying letter rotations is that leading neural models of visual word recognition make explicit predictions about how letter rotations affect word processing. The SERIOL model of word recognition (Whitney, 2001, 2002) assumes that “input levels to letter units are reduced for rotated stimuli” (p. 119, Whitney, 2002), increasing gradually the time required to encode a letter string. Instead, the Local Combination Detectors model (henceforth, LCD model; Dehaene et al., 2005) offers an all-or-none account on how letter rotation affects letter/word processing. Specifically, the LCD model postulates that letter rotation impedes the normal mapping of the letter features from the visual word form onto the letter detectors for angles above 40°-45°, but

that there would be little cost for rotations below that boundary¹⁴ (see also Cohen et al., 2008; Vinckier et al., 2006). Consistent with this latter prediction, a number of experiments have shown a cost in processing when words are rotated as a whole at 45° of rotation or larger. For instance, Cohen et al. (2008) found no differences in response times to words presented with 0° and 22.5° rotation angles, whereas there was a dramatic slow-down at angles of 45° and larger. In addition, other experiments have reported slower word identification times for rotated words at angles of 45° or larger than for their horizontally-presented counterparts (e.g., Barnhart & Goldinger, 2013 [0° vs. 90°]; Gomez & Perea, 2014 [0° vs 45° vs. 90°]; Koriat & Norman, 1984 [0° vs. 60° vs. 120° vs. 180°]).

Clearly, the above-cited studies consistently showed a cost when reading rotated words at 45° or larger. However, they were not informative as to the locus of the effect. These effects could have occurred at an early encoding stage—as the LCD or SERIOL models would predict—but they could also have been due to a processing cost at a late verification stage of processing where the visual percept is checked against the stored lexical unit. To directly test whether the locus of the effect of rotation occurs early in word processing, Perea et al. (2018) conducted a masked priming lexical decision experiment with 90° rotation of the whole word. The rationale of their experiment was that if the encoding of orthographic-lexical representations were hindered by rotations of 45° or larger, one would expect a dramatically reduced masked priming effect for 90° rotated words. However, they found sizeable masked identity and transposed-letter priming effects in line with the effects obtained with horizontally presented stimuli (see also Yang & Lupker, 2019, for evidence of robust masked priming effects with 90° and 180° rotated words). Hence, at least when words are presented in isolation, word rotation does not seem to hinder initial access to the orthographic representations (see also Perea et al., 2020, for similar evidence with isolated letters).

¹⁴ Note that the prediction of the LCD model (Dehaene et al., 2005) “letter detectors should be disrupted by rotation (>40°)” (p. 340) was based on a study conducted by Logothetis and Pauls (1995) with monkeys. Nevertheless, the empirical evidence provided by Cohen et al. (2008) and Vinckier et al. (2006) with adult readers was interpreted in light of this prediction of the model.

The robustness of masked priming effects to word rotations is quite revealing. However, one could argue that readers might rotate the whole stimulus and then process each letter as usual (see Gomez & Perea, 2014; Whitney, 2002, for discussion). Thus, a more stringent test for the LCD and SERIOL models would be to examine the effect of rotating the individual letters within words rather than rotating the whole word. Kim and Straková (2012) conducted a lexical decision experiment in which individual letters of words and pseudowords were rotated (0° , 22.5° , 45° , 67.5° , or 90°). Importantly, they also recorded event-related potentials (ERPs) during the task. Kim and Straková (2012) found an increase in amplitude in the P1 and N170 components (i.e., two indexes that reflect the initial contact with word-form features) for all rotation angles (i.e., including 22.5°) when compared to the horizontal format. To explain this pattern, they suggested that words whose individual letters are rotated require additional resources that are recruited very early in processing. The behavioral data showed a different picture, though. For words, response times were remarkably similar for rotation angles between 0° - 67.5° , with a large increase from 67.5° to 90° . Accuracy rates were similar for angles from 0° to 45° , slightly decreased from 45° to 67.5° , and there was a large drop from 67.5° to 90° . Thus, the behavioral data from Kim and Straková (2012) suggest that word recognition was not severely affected by rotations smaller than 45° - 67.5° —note that evidence for early ERP effects does not necessarily imply longer word identification times (see Perea et al., 2015, for discussion).

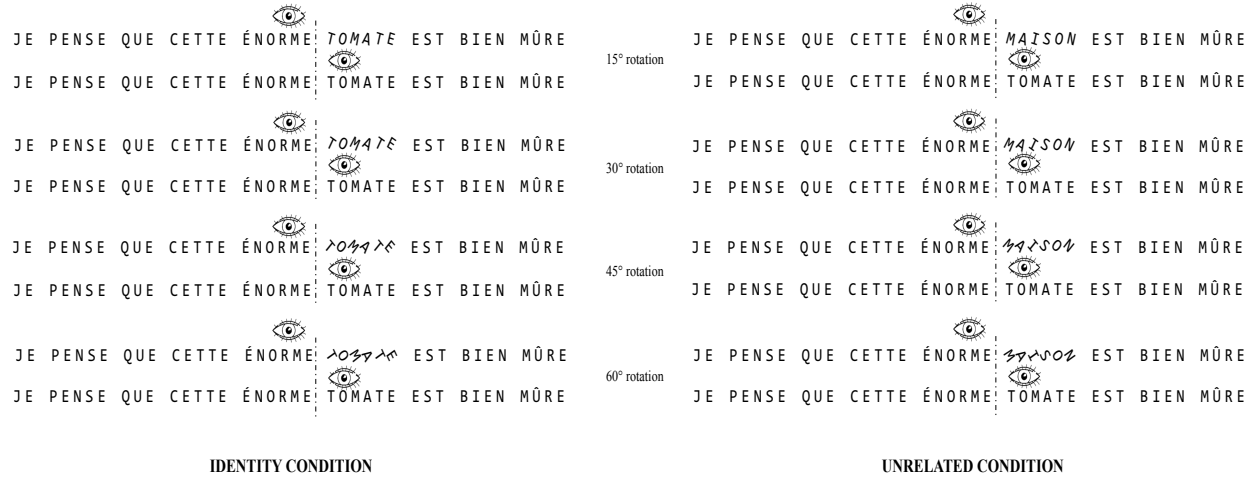
More important for the present purposes, a recent study measured participants' eye movements when reading sentences in which the individual letters within words were rotated (30° , 60°) or not (Blythe et al., 2019). Blythe et al. found an overall reading cost on sentence reading (total viewing time, number of fixations) for both rotation angles (30° and 60°) when comparing to upright presented text and this cost was larger at the 60° than at the 30° rotation angle. They also examined whether rotation angle interacted with word frequency by embedding a high- vs. low-frequency target word in each sentence. While both rotation angle and word-frequency had additive effects in the first-fixation duration on the target word ($30^\circ < 60^\circ$; high-frequency < low-frequency), total fixation

durations (i.e., the sum of all fixations in the target word) showed an interaction between the two factors: The magnitude of the word-frequency effect was a function of the rotation angle (i.e., the difference in eye fixation durations between high- and low-frequency target words was greater at 60°, it decreased at 30° and it was even smaller at 0°), thus suggesting that lexical access was impaired even by relatively small rotations (i.e., 30°). Blythe et al. (2019) concluded that letter detectors are affected by rotations well below 45° during sentence reading, and thereby the LCD predictions concerning letter rotations should be revised (i.e., the gradual cost of letter rotation may fit better with the SERIOL model). Although Blythe et al.'s (2019) experiment confirmed that letter rotations do affect the reading process, it does not inform us as to whether the effect of letter rotation occurs early in processing (i.e., as posited by the LCD and SERIOL models) or whether it occurs at late integration processes (keep in mind that participants could have adjusted their eye movement pattern to the rotation angle of the sentences).

The main aim of the present experiment was to examine whether the cost due to letter rotation in normal silent reading emerges early in processing, as predicted by the LCD and SERIOL models. To obtain a strict test of the resilience of letter detectors to rotations, we rotated the individual letters within words in a parafoveal word during sentence reading with the gaze-contingent boundary change paradigm (Rayner, 1975). Rayner's gaze-contingent boundary change technique allows for the manipulation of parafoveal information that is available to the reader before their eyes move to that location hence enabling foveal processing of a target word presented at the same location (see Figure 1). The rationale of this paradigm is that when we read, we extract information not only from the fixated word, but also from the following word(s) in the parafovea (Rayner et al., 1981, 1982; see also Vasilev et al., 2018, for evidence with degraded previews). Thus, it is possible to infer the amount of useful information the reader obtains from the parafovea by examining how the fixation durations on a target word vary as a function of the type of parafoveal preview (e.g., identity previews vs. unrelated previews; phonological preview vs. orthographic control preview, etc.; see Rayner et al., 2012, for review).

Figure 1

Description of an eye movement contingent display-change trial with the eight experimental conditions



Note. Identity preview rotated at 15°, 30°, 45° and 60°; “*I think this huge tomato is ripe*”; unrelated preview rotated at 15°, 30°, 45° and 60°; “*I think this huge house is ripe*”. The eye symbol represents where the reader is fixating, and the vertical dashed line represents the invisible boundary preceding the target word. Before crossing the boundary, the sentence is presented with the rotated previews. When the eyes cross the boundary, the parafoveal preview is replaced by the target word in the canonical format.

In the present gaze-contingent boundary change paradigm experiment, the parafoveal preview could be the same as the target (identity preview-target pairs) or a different word (unrelated preview-target pairs). We chose identity previews because they provide the largest effects in the parafovea (see Angele et al., 2013). The letters of the identity/unrelated previews could be rotated at four different angles, the difference always being 15° (15°, 30°, 45° and 60°; see Figure 1). Critically, the letters of the fixated target word—as well as the rest of the sentence—were always in the canonical upright orientation. We employed 15° of rotation as a baseline rather than 0° to avoid a perceptual continuity between preview and target in the identity condition.¹⁵ All the sentences were presented in UPPERCASE, the reasons being that: 1) in lowercase, some letters could be ambiguous when rotated (e.g., the rotated letter b can be confused with the letters p or q); 2) in lowercase words, low spatial frequency information

¹⁵ Keep in mind that any difference between a 15° and the canonical condition could be due to the congruency of format (both upright) that would induce the 0° rotation angle.

spanning the whole word (e.g., compare *bishop* vs. *BISHOP*) would have been less disrupted for smaller than for larger rotation angles (i.e., the whole visual word form of lowercase words could be more informative for smaller than for larger rotation angles). Of note, the pattern of eye movements when reading sentences is similar when written in lowercase or uppercase (see Perea et al., 2017).

The predictions were the following. First, if, as hypothesized by the LCD model (Dehaene et al., 2005), letter detectors are hindered by rotations in the first stages of orthographic-lexical processing only when they are above 40°-45°: 1) we would expect shorter fixation durations on the target word after an identity preview than after an unrelated preview, and to a similar degree, at 15° and 30° rotations (i.e., well below 45°); and 2) we would expect a dramatic drop in preview effects when the letters of the parafoveal preview are rotated 45° or 60°. Second, if, as hypothesized by the SERIOL model (Whitney, 2002), letter rotation gradually decreases the activation reached by the letter nodes, one would expect progressively smaller advantages of the identity preview as a function of rotation angle. A final scenario is that the orthographic coding system is widely tuned to rapidly identify rotated and/or distorted letters with little cost (see Hannagan et al., 2012; Yang & Lupker, 2019, and Perea et al., 2018, 2020)—in this case, we would expect an advantage of the identity condition over the unrelated condition not only for the smaller rotation angles (15° and 30°), but also for the larger rotation angles (45° and 60°).

Method

Participants

Forty participants (18 female) were recruited at Aix-Marseille University. The participants were all native French speakers and received either course credit or monetary compensation (€10/h). All participants reported normal or corrected-to-normal vision and ranged in age from 18 to 38 years old ($M=21.1$ y.o., $SD=4.11$). They were naïve to the purpose of the experiment and signed an

informed consent form before starting the experiment. Ethics approval was obtained from the Comité de Protection des Personnes SUD-EST IV (No. 17/051), and this research was carried out in accordance with the provisions of the World Medical Association Declaration of Helsinki.

Apparatus

Stimuli were displayed using OpenSesame (Mathôt et al., 2012) with each sentence occupying a single line. Eye movements were recorded with an EyeLink 1000 system (SR Research, Mississauga, ON, Canada) with high spatial resolution (0.01°) and a sampling rate of 1000 Hz—this device only recorded the right eye movements. The sentences were presented on a 20-inch ViewSonic CRT monitor with a refresh rate of 150 Hz and a screen resolution of 1024×768 pixels (30×40 cm). Stimuli were presented in upper case 18-point monospaced font (droid sans mono; the default monospaced font in OpenSesame) and the text was presented in black (0.15 cd/m^2) on a white background (21.70 cd/m^2); crosses and dots were presented with the same colors. Participants were seated 86 cm from the monitor, such that every 3 characters equaled approximately 1° of visual angle. The maximum validation error accepted was $.63^\circ$ of visual angle. A chin-rest and forehead-rest were used to minimize head movements.

Materials

We created 250 sentences in French, each contained an average of nine words (range 7-10). Sentences were created so as to minimize semantic predictability of the target words. This was verified via a cloze task in which the initial part of each sentence—until the word preceding the target word—was presented to four naïve individuals that were asked to predict the following word—none of these individuals took part in the experiment. The percentage of words that was predicted from the previous context was very low ($< 1\%$). The pre-target and target words had an average length of seven letters and an average Zipf frequency of 5.71 and 6.13, respectively (van Heuven et al., 2014). Word frequencies were obtained using the film subtitle frequencies of the Lexique2 database (New et al., 2004). For each target word, we created eight parafoveal

previews: identity condition with i) 15° of rotation; ii) 30° of rotation; iii) 45° of rotation; iv) 60° of rotation; and unrelated condition with v) 15° of rotation; vi) 30° of rotation; vii) 45° of rotation; and viii) 60° of rotation (see Figure 1). The parafoveal previews were counterbalanced across eight lists following a Latin square design. Each participant received 30 trials in each of the eight conditions and the sentences were presented in a different random order. The complete set of sentences, including the parafoveal previews, is presented in the Appendix.

Procedure

The experimental session took place in a dimly lit room. Participants were first informed about the experiment procedure. Their task was to read sentences for comprehension on a computer screen while their eye movements were registered. They were told that there would be comprehension questions after each sentence. The experiment procedure began with a five-point calibration phase in which participants had to look at individual dots on the screen. This was followed by 10 practice sentences to familiarize them with the procedure. Each trial had the following arrangement. First, a drift correction dot located 12 pixels (to the right of the left edge) was displayed. Second, the sentence was presented once the participant steadily looked at the drift correction dot—the starting point of the sentence was located just to the right of the dot. The display change occurred when the participant's eyes crossed an invisible boundary located between the target word and the word before (see Figure 1). Another invisible boundary was defined close to the end of the sentence (50 pixels before), such that the sentence disappeared when the eyes crossed that boundary (i.e., when participants read the last word). Finally, participants were shown a yes/no comprehension question after each trial that allowed us to verify that they had read for comprehension—participants were asked to press a corresponding key on a gamepad.

Data analyses

The critical idea underlying Rayner's (1975) boundary technique is that the parafoveal preview allows some preprocessing of the target word. Specifically, if the reader extracts useful information from the rotated-identity parafoveal preview

than for the rotated-unrelated parafoveal preview, the processing time on the target word would be reduced when directly fixated, thus resulting in shorter fixation durations (see Rayner et al., 2012). We examined three early eye fixation measures on the target word (i.e., the critical region): 1) the duration of the initial fixation on the target word (FFD); 2) the duration of a fixation on the target word when it was the only fixation on that word in first pass reading (SFD) and 3) the sum of fixation durations before leaving the target word (GD).

For the inferential analyses, we analyzed the eye fixation data with linear mixed effects (LME) and Bayes Factors. We employed the *lme4* (Bates et al., 2015, 2019), *lmerTest* (Kuznetsova et al., 2016), *car* (Fox & Weisberg, 2019) and the *BayesFactor* 0.9.12-4.2 (Morey & Rouder, 2018) packages in R (R Core Team, 2020). The models included two fixed factors: (1) preview-target relationship and (2) preview rotation angle—each factor was zero-centered. Subjects and items were the random factors in the models. Variables were log-transformed to reduce highly positive skew of fixation durations and keep the normality assumption of LME models. Regarding LME analyses, the maximal random structure model that converged for the three dependent variables was (Dependent Variable) $\sim$ Relationship*Angle + (1+Relationship|Subject) + (1+Relationship|Item). The p -values for each main factor and the two-way interaction were calculated using Wald χ^2 tests. In case of a significant interaction, we computed the effect of preview-target relationship for each rotation angle with *emmeans* (Lenth et al., 2020).

Results

Participants were quite accurate in responding to the comprehension questions (mean accuracy was .91; SD = .0023). Trials where the display change was triggered either early or late were removed from the analyses (less than 1%). The raw data were pre-processed using EyeLink algorithms to detect saccades, fixations and eye-blinks—the deletion of eye-blinks resulted in 0.5% of the data lost in the target region. Trials where the target word was skipped in first pass

were excluded from the analysis (9.78%). Fixation durations of less than 80 ms and more than 800 ms and gaze durations longer than 2000 ms on the target word were removed from the analyses. Furthermore, we removed any fixation duration measures that were more than 2.5 standard deviation for measure, condition and subject as they are unlikely to represent normal reading—we removed the 2.07% of FFD observations, the 3.23% of GD observations and the 1.95% of SFD observations. The average eye fixation measures per condition across participants are presented in Table 1.

Table 1.

Means (in ms) and Standard errors (in brackets) as a Function of rotation angle and preview-target relationship.

<u>Parafoveal Preview</u>						
Rotatio n Angle	Preview- Target Relationship	First Fixation Duration	Single Fixation Duration	Gaze Duration		
15°	Identity	240 (2)	243 (3)	303 (5)	21	26
	Unrelated	262 (2)	269 (3)	322 (4)		
30°	Identity	252 (2)	258 (2)	302 (4)	8	13
	Unrelated	260 (2)	271 (3)	314 (4)		
45°	Identity	254 (2)	261 (3)	307 (4)	5	7
	Unrelated	259 (2)	268 (3)	317 (4)		
60°	Identity	256 (2)	265 (3)	306 (4)	4	3
	Unrelated	260 (2)	268 (3)	315 (4)		

Note: Values in **bold** represent the significant parafoveal preview effects.

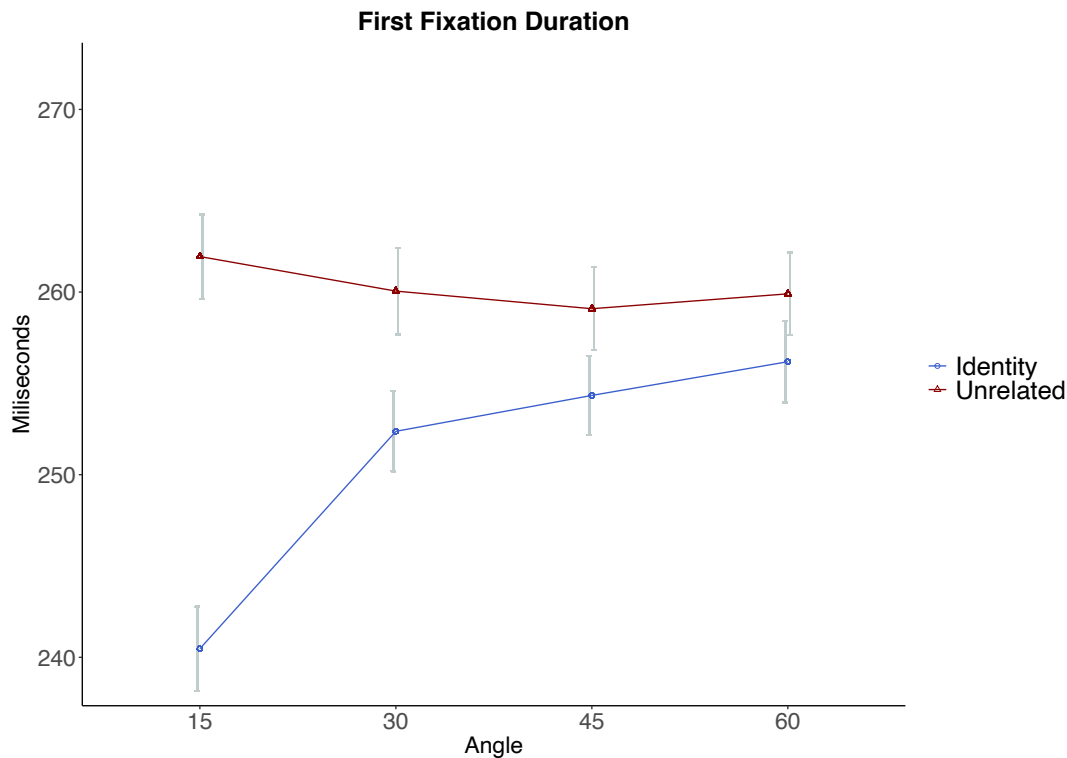
First Fixation Duration

On average, first-fixation durations on the target word were 9 ms shorter in the identity than in the unrelated condition ($\chi^2(1) = 25.616$, $p < .001$). We also found a main effect of angle, $\chi^2(3) = 14.840$, $p < .001$. More important, the interaction

between preview-target relationship and rotation angle was significant, $\chi^2(3) = 34.392$, $p < .001$. As can be seen in Figure 2, this interaction reflected that first fixation durations on the target word increased as a function of rotation in the identity condition, as reflected by significant linear ($t = 6.029$, $p < .001$) and quadratic ($t = -3.02$, $p = .002$) contrasts. For unrelated pairs, the duration of first fixations did not change as a function of rotation (none of the contrasts approached significance, all $|t| < 1.25$, $ps > .21$). The interplay between the two factors (i.e., preview-target relationship and angle) resulted in shorter first-fixation durations on the target word in the identity condition than in the unrelated condition when the letters of the preview were rotated 15° (21 ms; $b = -.089$, $SE = .012$, $t = -7.627$, $p < .001$) and, to a lesser degree, when rotated 30° (8 ms; $b = -.028$, $SE = .012$, $t = -2.418$, $p = .016$). Finally, the advantage of the identity condition did not reach significance when the letters of the preview were rotated 45° (5 ms) or 60° (4 ms) (both $ts < -1.146$, $ps > .253$). To examine whether the null effects for the 45° and 60° was due to *evidence of absence* rather *absence of evidence*, we computed Bayes Factors, as these offer valuable information to express our degree of certainty that the rotation angle did not affect the effect of preview-target relationship. For computing the Bayes Factors, we employed the `lmBF` function from the *BayesFactor* package with the default Cauchy distribution (centered around 0 and with a width parameter $\delta = 0.707$) (see Rouder et al., 2009; Wagenmakers et al., 2017, for discussion). The Bayes Factors for the 45° and 60° rotations were $BF_{01} > 7.508$, thus providing strong evidence toward the null hypothesis.

Figure 2

Mean First Fixation Durations by Condition (Identity vs. Unrelated) and Angle (15°, 30°, 45°, and 60°)



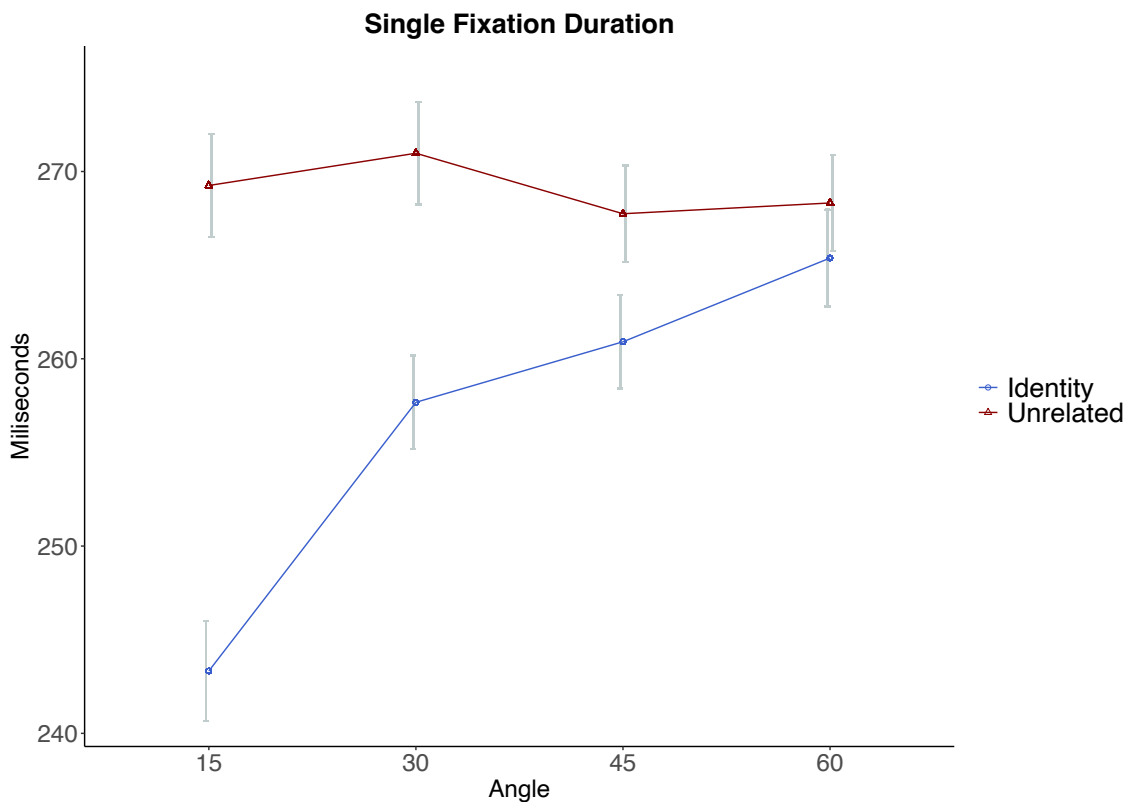
Single Fixation Duration

The proportion of first-pass trials with a single fixation in the target word was .66. Single fixation durations were, on average, 12 ms shorter in the identity than in the unrelated condition ($\chi^2(1) = 43.875$, $p < .001$). The effect of angle was also significant ($\chi^2(3) = 30.230$, $p < .001$) as was the interaction between the two factors ($\chi^2(3) = 41.251$, $p < .001$). This interaction showed that single fixation durations on the target increased non-linearly with the eccentricity of letter rotation in the identity condition—this was reflected by significant linear ($t = 8.414$, $p < .001$) and quadratic ($t = -3.205$, $p = .001$) contrasts (see Figure 3). The duration of single fixations for unrelated preview-target pairs remained constant independently of the rotation angle (none of the contrasts approached significance, all $|ts| < .76$, $ps > .45$). As a result of this interaction between preview-target relationship and angle, we found that single fixation durations were shorter in the identity condition than in the unrelated condition when the

preview was rotated 15° (26 ms; $b = -.111$, $SE = .012$, $t = -9.357$, $p < .001$) and, to a lesser extent, 30° (13 ms; $b = -.045$, $SE = .012$, $t = -3.831$, $p < .001$). For the 45° rotation angle, the 7-ms advantage in the identity condition, approached significance ($b = -.014$, $SE = .012$, $t = -1.187$, $p = .079$), but the evidence favored, if anything, the null hypothesis, $BF_{01} = 1.742$. Finally, for the 60° rotation angle, we found strong evidence toward the null hypothesis (3-ms advantage in the identity condition; $b = -.014$, $SE = .012$, $t = -1.187$, $p = .236$; $BF_{01} = 6$

Figure 3

Mean Single Fixation Durations by Condition (Identity vs. Unrelated) and Angle (15°, 30°, 45°, and 60°)



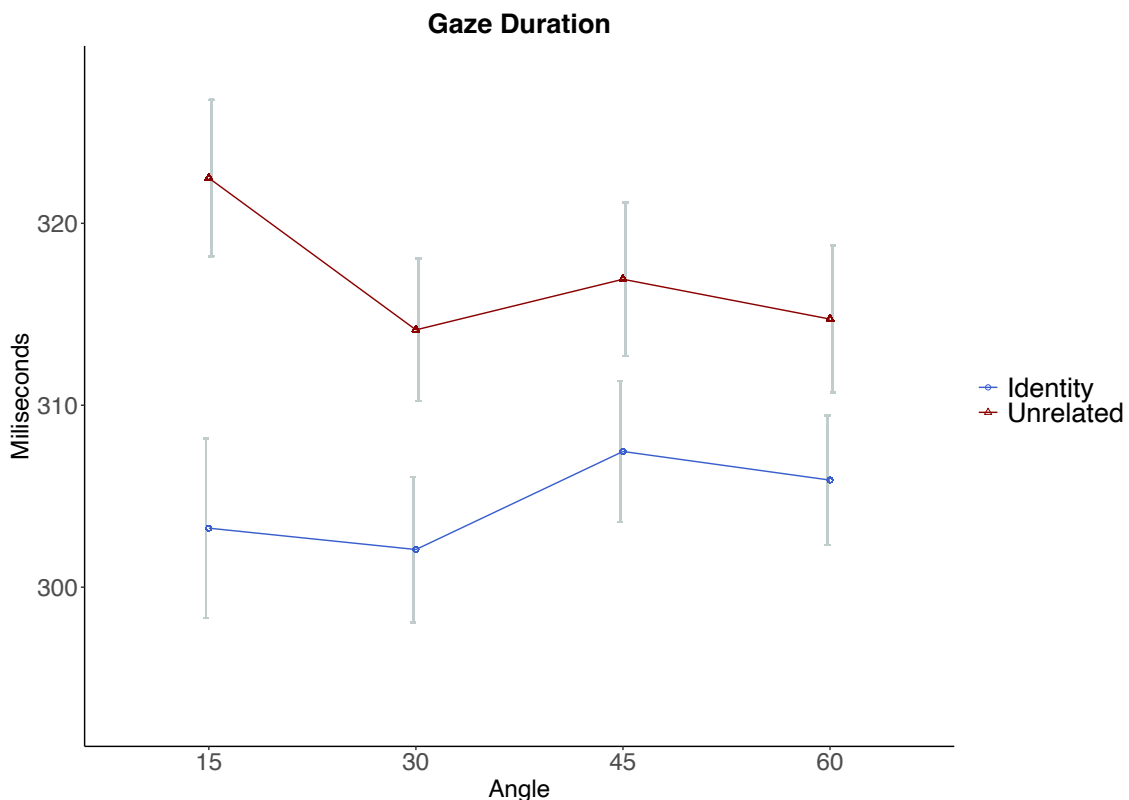
Gaze Duration

On average, gaze durations on the target word were 12 ms shorter in the identity than in the unrelated condition ($\chi^2(1) = 41.022$, $p < .001$). The effect of angle did not approach significance ($\chi^2(3) = 5.339$, $p = .149$) but, more importantly, the interaction between preview-target relationship and angle was significant ($\chi^2(3) = 20.258$, $p < .001$). This interaction showed that gaze durations on the target

word increased linearly with the eccentricity of letter rotation in the identity condition—this was reflected by significant linear contrasts ($t = 8.414$, $p < .001$). The duration of gaze durations for unrelated preview-target pairs remained constant independently of the rotation angle (none of the contrasts approached significance, all $|ts| < 1.52$, $ps > .12$). As a result of this interaction between preview-target relationship and angle, we found a 19-ms advantage of the identity condition 15° rotation ($b = -.093$, $SE = .013$, $t = -6.985$, $p < .001$), which was reduced to 12 ms at the 30° rotation ($b = -.052$, $SE = .013$, $t = -3.863$, $p < .001$). For the 45° rotation, we found a small 10-ms advantage of the identity condition ($b = -.027$, $SE = .014$, $t = -2.000$, $p = .046$), which actually slightly favored the null hypothesis, $BF_{01} = 1.688$, whereas for the 60° rotation, we found strong evidence for the null hypothesis (9 ms; $b = -.019$, $SE = .014$, $t = -1.398$, $p = .163$; $BF_{01} = 6.134$) (see Figure 4).

Figure 4

Mean Gaze Durations by Condition (Identity vs. Unrelated) and Angle (15°, 30°, 45°, and 60°)



Discussion

The present experiment was designed to test whether letter detectors are severely hampered by rotations above 40-45° during sentence reading, as postulated by the LCD model (Dehaene et al., 2005) or, whether this cost would emerge gradually as a function of increasing letter rotation, as hypothesized by the SERIOL model (Whitney, 2002). To that end, we conducted a reading experiment with sentences in uppercase using Rayner's (1975) gaze-contingent boundary change paradigm in which we examined whether identity parafoveal previews were more effective than unrelated parafoveal previews when the previews were rotated at 15°, 30°, 45°, or 60° (see Figure 1 for illustration). We found a remarkably similar pattern of findings for all three eye fixation measures on the target word (i.e., first-fixation duration, single fixation duration, and gaze duration). Fixation durations were shorter when preceded by an identity preview than when preceded by an unrelated preview and, critically, this difference was modulated by rotation angle: the advantage of the identity preview condition was greatest when the letters of the preview were rotated 15° and it was numerically smaller but still robust at 30° (see Table 1). Instead, for the 45° rotations the advantage of the identity preview condition was only marginal, and the Bayes Factor analysis did not produce clear evidence in either direction—in particular for single-fixation durations and gaze durations. Finally, for the 60° rotations, there were no signs of an advantage of the identity preview condition—indeed, the evidence in favor of the null effect was strong (all $BF_{01} > 6.134$). Thus, the present results extend and qualify the findings obtained by Blythe et al. (2019) with a technique that specifically focused on early word recognition processes during sentence reading.

This pattern of results poses some problems to the predictions of the LCD model concerning the impact of letter rotation on word recognition (Dehaene et al., 2005). Although letter detectors appear to be severely disrupted at rotation angles of 45° or larger in the very early moments of processing, letter rotation also has an effect for rotation angles well below 45° in the parafovea. In the three eye fixation measures, we found a greater identity preview advantage for 15°

than for 30° rotated previews. These findings, together with the data from Blythe et al. (2019) with 30° rotations, suggest that the disruption caused by letter rotations is not due to a sudden all-or-none shift at around the 40-45° boundary posited by the LCD model. Thus, the more parsimonious account of the current findings is that the amount of visual input reaching the letter detectors would decrease as a function of letter rotation, as proposed by the SERIOL model (see Whitney, 2002). This is also consistent with the increase in amplitude in ERP components that reflect the initial contact with word-form features (e.g., N170) even for small rotations (22.5°) of individual letters in words (Kim & Straková, 2012). Second, while the evidence of an identity preview advantage for 45° rotations was marginal, we found strong evidence for a null effect for 60° rotations, thus again suggesting a gradual nature of the cost due to letter rotations on the letter detectors of the word recognition system. Taken together, this pattern would suggest that the identity preview advantage of rotated previews during sentence reading is structured by rotation angle (i.e., greatest at 15°, intermediate at 30°, marginal at 45°, and absent at 60°). Further experiments should be conducted to test whether this pattern of results also occurs with other manipulations (e.g., phonologically or graphemically preview-target related words).

Importantly, one might argue that our findings are at odds with the presence of robust masked priming effects with rotated words found by Perea et al. (2018) and Yang and Lupker (2019). However, there are some key methodological differences between the present and the just mentioned studies. First, in the Perea et al. (2018) and Yang and Lupker (2019) experiments, the whole word was rotated, whereas in the present experiment we rotated the individual letters of the word (i.e., the manipulation was more extreme in the present experiment). Second, in above-cited masked priming experiments: 1) both primes and targets were presented foveally instead of parafoveal-then-foveal as in the present study (and the quality of visual information is higher in the fovea); and 2) both primes and the targets were presented rotated at the same angle. In the present experiment, the previews were presented parafoveally at different rotation angles, but the target word (as well as the rest of the sentence)

was always presented in the canonical upright format. Thus, our participants were reading normally oriented sentences, thus avoiding the adoption of specific strategies that might arise when reading oriented text. To further elucidate these apparent discrepancies, future research could examine the time course of masked priming words with individually rotated letters (e.g., using electrophysiological measures).

Our findings also have important implications for models of eye movement control in reading. At present, no current models of eye movement control during reading have explicit predictions on the processing of words composed of rotated letters (e.g., E-Z Reader, Reichle et al., 1998; SWIFT, Engbert et al., 2005; OB1-reader, Snell et al., 2018). This is because these models have focused on key elements of eye movement control (e.g., *when* and *where* to fixate the eyes) rather than on visual/sublexical/lexical processes. Our findings highlight the necessity of expanding these models to include early processes, including the mapping of visual information onto sublexical orthographic or phonological stages to offer a comprehensive picture of the reading process. As an example, the E-Z Reader model (Reichle et al., 1998, 2003) assumes that word identification is completed in two stages, reflecting early and late lexical processing. The first stage corresponds to the beginning of the identification of the orthographic form of the word, whereas full lexical access is reached in the second stage. The identity preview advantage observed in fixation durations when the letters of the rotated preview were 15°-30° suggests that readers can complete the first stage of lexical processing (i.e., extract abstract letter codes) even when the individual letters of words are rotated 30° (i.e., readers are able to extract enough information to begin programming a saccade to the next word even when that information is extracted from a word with letters rotated 30°) (see Angele et al., 2016, for a proposal of the mechanisms at play in parafoveal processing). As stated earlier, the cost of processing these rotated letters would be gradual (i.e., the larger the rotation angle, then smaller the information acquired in the parafovea).

In sum, skilled readers can efficiently convert parafoveal words composed of 15°-30° rotated letters to a stable orthographic code during silent sentence

reading, but not when the rotation angles are 45° or above. Thus, our data are in accordance with a gradual cost of letter rotations in the uptake of parafoveal information as a function of rotation angle, as posited by the SERIOL model of word recognition (Whitney, 2001, 2002). These findings should encourage the implementation of a more detailed account of the interface between visual form information and orthographic-lexical representations in computational models of eye movement control during sentence reading.

On the time course of the resilience of letter detectors to rotations: A masked priming ERP investigation

Abstract

A straightforward prediction of the Local Combination Detectors [LCD] model of word recognition (Dehaene et al., 2005) is that letter rotations above 40-45° should disrupt the mapping of the visual input onto orthographic representations. However, the evidence supporting this claim is scarce and not conclusive. To shed light on this issue, we conducted a masked repetition priming lexical decision experiment while recording the participants' EEG measures. Targets were always presented in the standard horizontal format, and we rotated the individual letters of the identity/unrelated primes (0°, 45°, or 90°). Behavioral and Event-Related Potentials (ERP) results revealed that the identity priming effect decreased as a function of letter rotation. Importantly, the ERP data allowed us to examine in detail the time course of processing of words with rotated letters. Amplitude comparisons showed that identity priming followed the typical course for 0° primes (i.e., it started around 100 ms, in the visual feature encoding stage, and strengthened with processing time). The parallel effect for 45° primes emerged later, at around 175 ms. This pattern strongly suggests that letter rotations at around 45° have a processing cost, thus providing evidence in favor of the LCD model of word recognition (Dehaene et al., 2005).

Key words: letter rotation, word processing, masked priming, ERPs.

Introduction

Learning to read allows us to become autonomous and self-developing individuals in today's society. On a big scale, becoming literate individuals transforms our knowledge and views about the world. On a smaller personal scale, it changes the very subtle cognitive and brain mechanisms that connect visual forms with meaning. Ultimately, reading acquisition leads to building a flexible and powerful decoding system: despite the multiple variations in the visual form of written words (e.g., house, *house*, *house*, *house*, *house*, etc.), adult skilled readers can rapidly access their corresponding lexical-semantic representations (see Rayner et al., 2012). Indeed, even briefly shown prime stimuli with heavily distorted letters (e.g., CATCHAPs like *house*) produce sizeable masked identity priming effects (see Hannagan et al., 2013). The enormous efficiency of lexical access in adult readers demonstrates that the word recognition system has developed some tolerance to noise and variations in the visual form of the words' constituent letters (see Cohen & Dehaene, 2009; Grainger & Dufau, 2012). However, the empirical evidence on the resilience of the word recognition system to perceptual distortion is still scarce, probably because it entails many different elements (see Vergara-Martínez et al., 2021, for discussion).

The main goal of the present study is to examine to what degree the visual word recognition system is resilient to the degradation in the visual form of the letters. Given the drawbacks of investigating the distortion of the visual form of written text due to the many potential factors at play, we focused on a single parameter: letter rotation. Crucially, the examination of the impact of letter rotation in word recognition has an added value: a prominent neural-inspired model of visual word recognition (e.g., Local Combination Detectors [LCD] model, Dehaene et al., 2005) makes an explicit prediction on the influence of rotations in visual word processing: "letter detectors should be disrupted by rotation ($> 40^\circ$)" (Dehaene et al., 2005, p. 340). Specifically, the LCD model assumes that the visual word recognition system shifts from the parallel encoding of letters to

an effortful serial processing strategy for rotation angles above 40°- 45°, thus hampering lexical processing (see Cohen et al., 2008)¹⁶.

The present research aimed to test this prediction of the LCD model. To that end, we conducted a masked priming lexical decision task in which the individual letters of the primes (that could be identical or unrelated to the target) were rotated in 0°, 45°, or 90°. Importantly, we also recorded the participants' brain electrical activity with an electroencephalogram (EEG) while doing the task. The EEG measures allowed us to characterize the impact of letter rotation during the time-course of lexical access. In the following lines, we first review the relevant background of the effects of letter word rotation effects on lexical processing. Next, we will introduce the experiment along with the predictions.

As indicated above, the tenets from the LCD model set limits to the tolerance of the visual word recognition system regarding letter rotation in 45°, meaning that letter rotations at around 45° or above would interfere with normal visual word recognition. In an event-related potential (ERP) lexical decision experiment, Kim and Straková (2012) examined the impact of letter rotation in visual word recognition. Participants were presented with words (and nonwords) where individual letters were rotated 0°, 22.5°, 45°, 67.5°, or 90°. Behavioral results showed that both response accuracy and speed decreased non-linearly with the eccentricity of letter rotation, with the greatest drop at the 67.6° – 90° rotation step. Electrophysiological results revealed early effects of letter rotation in the P1 and N170 ERP components, two ERPs that mark the beginning of word-form analysis during word recognition (i.e. P1 and N170 amplitudes are enhanced in response to letter strings relative to nonlinguistic stimuli; Bentin et al., 1999; Maurer et al., 2005; Rossion et al., 2003). Kim and Straková (2012) found that increasing letter-rotation led to monotonic increases in the P1 and N170 amplitudes up to 67.5°. However, at 90°, P1 and N170 amplitudes decreased substantially. To explain this pattern, they suggested that moderate deviation

¹⁶ This claim was initially inspired by a study on the generalization at recognizing isolated objects at different orientations in macaques (Logothetis & Pauls, 1995). Of note, Cohen et al. (2008) and Vinckier et al. (2006) interpreted their data with adult readers in light of this invariance-within-limits prediction of the model.

from preferred stimulus properties (less than 45°) can recruit additional processing resources. However, deviation beyond a certain level (e.g., letter rotations beyond 67.5°) exceeds the resiliency of the processing system, leading to a decrease in the EEG response.

Are these results consistent with LCD predictions? As noted by Kim and Straková (2012), the LCD cannot readily explain why letter rotation *gradually* enhance brain responses to words composed of rotated letters from 0 to 22.5° and 45°. Kim and Straková (2012) suggested that the additional processing resources underlying the enhanced brain responses for these letter rotations could imply top-down feedback (attention-driven amplification; see Qiao et al., 2010; Vergara-Martínez et al., 2021) to low-level form processing when the letters are presented in an unfamiliar format (i.e., rotated). That is, top-down attentional processes would help override the disruptive effect of rotation at an abstract orthographic processing level (see Carreiras et al., 2014). One limitation of the procedure used by Kim and Straková (2012) is that word response times were not collected online but after a signal 800 ms post-target. Thus, it is not clear whether their early ERP effects necessarily reflected a cost on lexical access (see Perea et al., 2015, for discussion).

A more stringent test of the LCD model regarding letter rotation is to block out the attention-driven amplification elicited by distorted stimuli presented in isolation (Qiao et al., 2010; Vergara-Martínez et al., 2021). To that end, we employed a technique that taps the very initial steps of lexical access, while also collecting word response times. Experimental paradigms such as the parafoveal preview technique or masked priming address this issue. Notably, several recent studies with rotated words, using Forster and Davis' (1984) masked priming technique, have questioned the tenet of the LCD model regarding the 45° threshold of the resilient of letter detectors to rotation (e.g., Benyhe & Csibri, 2021; Perea et al., 2018; Yang & Lupker, 2019). The rationale of these experiments was that if letter rotations above 45° slow down the access to the word's orthographic representations, masked identity rotated primes would no longer facilitate (or only to a minimal extent) target processing. However, Perea et al. (2018) found sizeable masked priming effects with 90° rotated words (both

identity priming and transposed-letter priming), similar to those with canonical presented stimuli. Likewise, Yang and Lupker (2019) reported reliable transposed-letter priming effects not only with 90° but also with 180° rotated stimuli. While these findings demonstrate that the word recognition system is resilient to word rotations, these experiments suffer from an interpretive issue: the prime and the target were rotated. Thus, one might argue that participants could have been accustomed to mentally rotating the stimuli to their canonical disposition and then processing them as usual (see Gomez & Perea, 2014; Whitney, 2002, for discussion).

An excellent strategy to minimize the above issue is to use a paradigm where participants have to process an upright presented stimulus preceded by an identity (or unrelated) stimulus (in various degrees of rotation) that is not consciously perceived. Benyhe and Csibri (2021) followed this strategy in a masked priming experiment in which participants were asked to read aloud target words in their canonical orientation (0°). The targets were preceded by identical or unrelated prime words that could be rotated in 0°, 30°, 60°, 90°, 120°, 150° and 180°. They found significant priming effects for rotation angles until 60° when the prime duration was 50 ms and 90° when the prime duration was 75 ms. Thus, masked priming evidence with word rotations favors the idea that letter detectors are resilient to rotations >45° even in the first moments of processing.

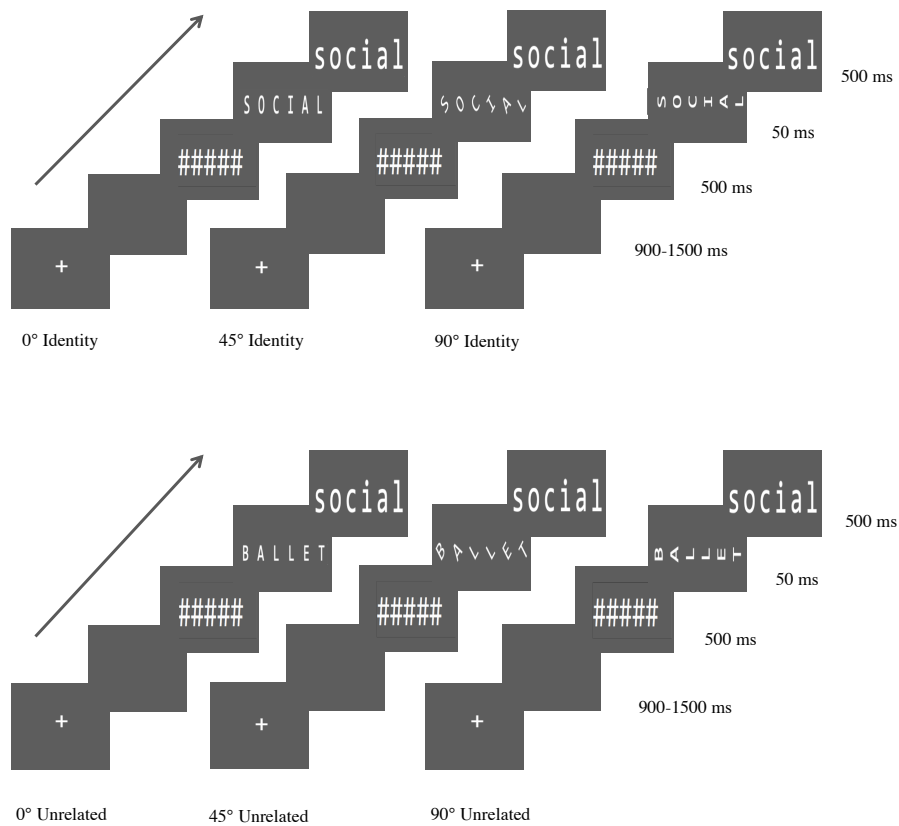
The robustness of masked priming effects to word rotations is quite revealing. However, in the above-cited experiments, one might argue that the cognitive system rotates the whole stimulus and then process each letter as usual (i.e., a mental rotation; see Gomez & Perea, 2014; Whitney, 2002, for discussion). This scenario is much less disrupting than rotating the letters that make up the words. Keep in mind that, unlike word rotations, the rotation on a letter-by-letter basis involves the disruption of trans-letter features (i.e., features that are larger than letters but smaller than words; Mayall & Humphreys, 1996). Thus, the effect of rotating the individual letters within words is a more stringent test for the invariance-within-limits assumption of the LCD model than the rotation of the whole words. This was the approach followed in a recent study by Fernández-López et al. (2021). They used Rayner's (1975) gaze-contingent boundary

change paradigm during sentence reading. The parafoveal preview was either identical or unrelated to the target word, and the letters of the preview were rotated 15°, 30°, 45°, or 60°. Notably, the letters of the fixated target word and the rest of the sentence were always in the canonical—upright—orientation. Results showed that the advantage of the identity preview condition in eye fixation times on the target word decreased progressively as a function of the rotation angle (i.e., the identity advantage was sizeable for 15° and 30°, weak for 45°, and absent for 60°). Notably, at the critical boundary of the LCD model (i.e., 45°), the advantage of the identity preview was only marginal (i.e., 5, 10, and 7 ms for first fixation duration, gaze duration, and single fixation duration, respectively). Thus, the cost of letter rotation of parafoveal previews during sentence reading increased as a function of rotation angle, being substantial after 45°, thus providing empirical support to the LCD model. One interpretive issue, however, is that in the Fernández-López et al. (2021) experiment, the prime (“preview” in the eye movement literature) was presented in the parafovea, where the quality of the spatial information is lower than in the fovea. Hence, it is difficult to completely ascertain whether the cost of letter rotations at 45° reflects the system bounds on early orthographic processing or whether it reflects structural limitations of parafoveal processing.

To thoroughly characterize the effects of individually rotated letters during word recognition and thus test the tenets of the LCD model, we chose the masked priming technique (Forster & Davis, 1984) in combination with the recording of EEG measures. The reason for selecting this paradigm is twofold: (i) it taps on the very early stages of visual word recognition, and (ii) it reflects only bottom-up processing within the ventral stream, in the absence of attention-driven amplification (Qiao et al., 2010). Furthermore, the invariance-within-limit assumption of the LCD model can be readily tested in masked priming experiments: letter rotations at 45° or higher would slow down orthographic processing, decreasing masked identity priming effects. As in the Benyhe and Csibri (2021) experiment, the targets in the present study were *always* presented in the canonical upright position so that participants were unaware of prime identity or prime rotation (Figure 1).

Figure 1

Depiction of the masked priming task



Thus, the main goal of the present experiment was to examine whether readers can extract abstract orthographic representations of rotated letters during visual-word recognition and, if so, when. Of specific interest was the scrutiny of three ERPs associated with the processing of visual feature representations (N/P150), abstract letter representations (N250), and lexical-semantic representations (N400) in masked priming experiments (see Holcomb & Grainger, 2006, 2007, for review). Starting approximately at 125 ms after the target onset and ending by 175 ms, the N/P150 reflects the initial mapping of the stimulus' visual features (e.g., size, color, cAsE) onto higher-level orthographic representations. Differences in the ERP waves in this epoch depend on the degree of overlap of coarse-grain visual features at a relatively abstract level, with no impact of lexical factors (e.g., lexicality or word frequency). For instance, the N/P150 component is modulated by letter case, regardless of whether the letters of prime and target match the same orthographic representation (e.g., it is larger

for altar-ALTAR than for ALTAR-ALTAR; Vergara-Martínez et al., 2015; see also Petit et al., 2006). The second ERP of interest peaks around 250 ms (N250; 175-300 ms) and reflects the mapping of the orthographic units onto orthographic word forms. The N250 component is modulated by the degree of orthographic overlap between prime and target: its amplitude is larger to completely unrelated pairs, moderate to almost-overlapped pairs, and most attenuated for identical pairs (e.g., agenda-SOCIAL > soical-SOCIAL > social-SOCIAL; see Chauncey et al., 2008; Dickson & Federmeier, 2014; Dufau et al., 2008; Pickering & Schweinberger, 2003, for further evidence). Finally, the N400 component, occurring approximately from 300 ms to 450 ms after stimulus onset, reflects the mapping of word representations onto semantic representations (i.e., lexical-semantic interactions) in long-term memory. Thus, similarly to the N250, the N400 component is more negative to targets that are completely unrelated to their primes, intermediate to targets that partially overlap their primes, and least negative to targets that are complete repetitions of their primes (e.g., agenda-SOCIAL > soical-SOCIAL > social-SOCIAL). Importantly, unlike the N250, the N400 in masked priming is sensitive to word stimuli (but not nonword stimuli) and has been associated the access to semantics (see Holcomb & Grainger, 2006; Kiyonaga et al., 2007).

In the present experiment, each target stimulus was preceded by a masked prime that could be the same as the target (e.g., social-SOCIAL) or unrelated (e.g., agenda-SOCIAL). Critically, the prime was composed of letters rotated 0°, 45°, or 90° (see Figure 1). If, as hypothesized by the LCD model (Dehaene et al., 2005), letter rotation above 40°–45° severely hinders the encoding of letter detectors in the first stages of orthographic-lexical processing, we would expect at the behavioral level: i) a sizeable identity priming effect for the canonical 0° primes; ii) a weak/negligible effect for 45° rotated primes, and 3) null priming effects for the more extreme 90° rotated primes. Alternatively, if the orthographic coding system is widely tuned to rapidly process distorted stimuli (Benyhe & Csibri, 2021; Hannagan et al., 2012; Perea et al., 2018, 2020; Yang & Lupker, 2019), one would expect a priming effect at all rotation angles. This

latter scenario does not exclude that priming effects could be numerically smaller for the steeper rotations (see Fernández-López et al., 2021).

More importantly, the registration of the participants' ERPs allowed us to track the size and time course of the masked identity priming effect for all rotation angles. The predictions in the framework of the LCD model are the following. We expect the N/P150 to be sensitive to coarse perceptual changes (i.e., letter rotation) between prime and target (see Gutierrez-Sigut et al., 2019; Vergara-Martínez et al., 2015). In the N250 component, we would expect identity priming effects for the 0° rotation primes. This difference would be substantially weaker for 45° rotated primes, as letter detectors would be at the boundary to automatically encode letter identities (Dehaene et al., 2005), and it would be negligible for 90° rotated primes. For the N400, we can envision two scenarios. On one side, as the N400 is sensitive to the same degree of orthographic overlap as the N250, one could expect the same pattern of results for both components (N250 and N400). On the other side, as the N400 reflects the access to the lexical entry of the target, top-down lexical feedback could override the disruption of prime letter rotation (see Carreiras et al., 2014; Vergara-Martínez et al., 2021). In this latter scenario, one can expect that top-down processes may facilitate the encoding of the rotated primes, thus leading to priming effects also with rotated primes—this may be more pronounced for the less disruptive 45° rotated primes.

The alternative hypothesis to the LCD predictions entails that letter detectors could quickly overcome the constraints imposed by rotation early in processing, similar to the rotation of whole words (Benyhe & Csibri, 2021; Perea et al., 2018, 2020; Yang & Lupker, 2019). In this case, one would expect sizeable identity priming effects in the N250 and N400 components at all rotation angles. This latter outcome would reflect that participants can rapidly integrate the abstract orthographic representations of the primes into the targets regardless of the rotation of the individual letters of the primes.

Method

Participants

A group of 26 (20 females) undergraduate university students participated in the study for course credit. Ages ranged from 19 to 30 years ($M = 21.57$, $SD = 3.42$). All participants were native Spanish speakers with no neurological or psychiatric impairment history and normal or corrected-to-normal vision. In addition, all participants were right-handed, as assessed with a Spanish version of the Edinburgh Handedness Inventory (Oldfield, 1971). This study was approved by the Experimental Research Ethics Committee of the University of València. Before starting the experiment, all participants signed an informed consent form.

Materials

We selected 240 target words of five and six letters from the EsPal subtitle-based database in Spanish (Duchon et al., 2013). The mean frequency was 122 occurrences per million (range: 0.5 – 1352)—this corresponded to an average Zipf frequency of 4.83 (range: 2.65—6.13; van Heuven et al., 2014). The mean orthographic Levenshtein distance 20 (OLD20; Yarkoni et al., 2008) was 1.50 (range: 1—2). We also created 240 nonwords as foils using Wuggy (Keuleers & Brysbaert, 2010). To act as primes, we employed the same stimuli selected for targets. Each target stimulus, presented in lowercase, was preceded by an uppercase prime that could be (a) the same as the target, or (b) an unrelated stimulus; the prime could also be rotated in (a) 0°, (b) 45° and (c) 90°. As in previous masked priming experiments, the unrelated primes for word targets were words and the unrelated primes for nonword foils were nonwords—note that the lexical status of the unrelated primes does not affect lexical processing (see Fernández-López et al., 2019). Importantly, in the typical scenario of masked priming lexical decision, the prime is presented in lowercase and the target is presented in uppercase. Nevertheless, to avoid binding and mix-up with the rotated letters (i.e., a rotated *b* can be confused with a *p* or a *q*), we rotated letters in uppercase words. To include the six priming conditions across all target words, we created six counterbalanced lists of materials in a Latin square manner (i.e.,

each target appeared once in each list, each time in a different priming condition). Different participants were assigned randomly to each list. To have the participants familiarized with the task, the 480-trial experimental phase (i.e., 40 items per condition) was preceded by a short practice composed of 12 trials (6-word trials and 6-nonword trials) with the same characteristics as the experimental trials. These practice trials were not included in the analyses.

Procedure

Participants sat comfortably in a dimly lit and sound-attenuated room. All stimuli were presented on a high-resolution monitor positioned slightly below the eye level, 85-90 cm in front of the participant. The size of the stimuli and distance from the screen allowed for a visual angle of fewer than 5 degrees horizontally. Stimuli were presented in white Consolas font against a dark-gray background. The primes were presented in uppercase and 20-pt of size, whereas both targets and the mask were presented in lowercase and 36-pt. Stimulus display was controlled by Presentation software (Neurobehavioral Systems).

In each trial, a pattern mask (i.e., a series of # signs that matched the length of the target item) was displayed for 500 ms in the center of the computer screen. An uppercase prime replaced the mask in the same spatial location for 50 ms, which was replaced by a lowercase target stimulus that remained in the screen for 500 ms. After participants' response (within an interval of 2 seconds following target presentation), there was a blank screen of random duration between 900 and 1,500 ms. Participants were asked to decide whether the target stimulus was a Spanish word or not by pressing the "yes" or the "no" key. Both speed and precision were stressed in the instructions. To minimize participant-generated artifacts in the EEG signal during the presentation of the experimental stimuli, participants were asked to refrain from blinking and moving from the onset of each trial to the set-up period after response. Instructions did not mention the existence of any uppercase stimuli. Every 30 trials, there was a brief pause for resting and impedance checking. The hand used for each response was balanced across participants. Lexical decision times were measured from target

onset until the participant's response. Participants did not receive feedback on reaction times or error rates during the experiment. The order of the trials was randomized for each participant. The whole session, including set up, lasted approximately 1 hour and 30 minutes.

EEG recording analysis

The electroencephalogram (EEG) was recorded from 29 Ag/ AgCl electrodes mounted in an elastic cap (EASYCAP GmbH, Herrsching, Germany). These electrodes were referenced to the right mastoid and re-referenced off-line to the averaged signal from two electrodes placed on the left and right mastoids. Eye movements and blinks were monitored with electrodes providing bipolar recordings of the horizontal and vertical (over the left eye) electrooculogram (EOG). The EEG recording was amplified and bandpass filtered between 0.01–100 Hz with a sampling rate of 250 Hz by a BrainAmp amplifier (Brain Products, GmbH, Gilching, Germany). An off-line bandpass filter between 0.01 and 20 Hz was applied to the EEG signal. Impedances were kept below 5 k Ω during the recording session. Epochs of the EEG corresponding to 200 ms pre- to 600 ms post-target onset were analyzed. Baseline correction was performed using the average EEG activity in the 200 ms preceding the onset of the target stimuli. Following baseline correction, trials with muscle activity, eye movements, or blink activity were rejected (8.95%). Trials with incorrect lexical decision responses (3.32%) were neither included in the average ERPs. All participants had a minimum of 30 acceptable correct trials per condition¹⁷ (ID-0°: M = 36, SD = 2.6; ID-45°: M = 35, SD = 2.3; ID-90°: M = 35, SD = 2.4; UN-0°: M = 35, SD = 2.3; UN-45°: M = 35, SD = 2.6; UN-90°: M = 35, SD = 2.6). ERPs were averaged separately for each of the experimental conditions, each of the subjects, and each of the electrode sites. We exclusively focused on the word trials because masked

¹⁷ To increase the signal-to-noise ratio in four subjects with a large number of eyeblinks, we applied the independent component analysis procedure to correct eye blinks using the Infomax algorithm (Jung et al., 2000).

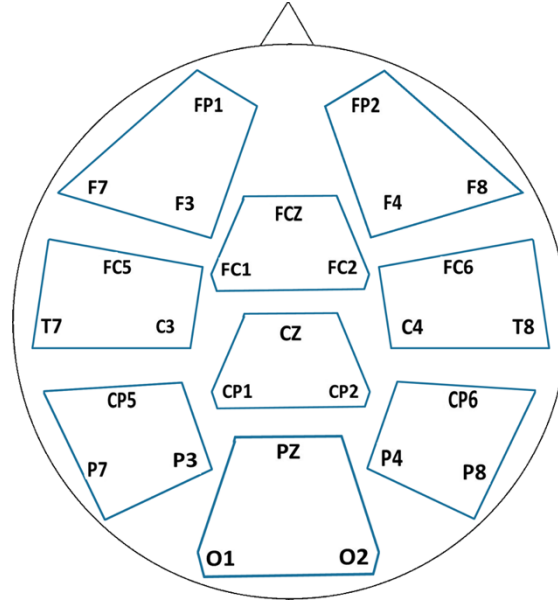
priming effects for nonword trials in the lexical decision task tend to be unreliable (see Carreiras et al., 2009; Gutiérrez-Sigut et al., 2019)¹⁸.

Statistical analyses were performed on the mean amplitude of three contiguous time windows (100-150 ms; 175-300 ms; 300-450 ms). This was done for the six experimental conditions defined by the combination of the factors Prime-Target Relation (identity, unrelated) and Rotation angle (0°, 45°, 90°). The selection of these epochs was motivated by our aim to track the time course of the potential differences between experimental conditions and was based on previous studies (Carreiras et al., 2007; Holcomb & Grainger, 2006; Grainger & Holcomb, 2009; Vergara-Martínez et al., 2015). Following a similar strategy in the related literature (e.g., Vergara-Martínez et al., 2020), we analyzed the topographical distribution of the ERP results by including the averaged amplitude values across three electrodes of nine representative scalp areas that result from the factorial combination of the Laterality (left, central, right) and Distribution (anterior, medial, posterior): left-anterior (FP1, F7, F3), left-medial (FC5, T7, C3), left-posterior (CP3, P7, P3), central-anterior (FZ, FC1, FC2), central-medial (CZ, CP1, CP2), central-posterior (PZ, O1, O2), right-anterior (FP2, F8, F4), right-medial (FC6, T8, C4); and right-posterior (CP6, P8, P4; Figure 2). For each time window, a separate repeated-measures analysis of variance (ANOVA) was performed, including the factors Laterality, Distribution, Rotation angle, and prime-target Relation. In all analyses, List (List 1, List 2, List 3, List 4, List 5, and List 6) was included as a between-subjects factor to extract the variance due to the counterbalanced lists (Pollatsek & Well, 1995). Critical values were adjusted using the Greenhouse and Geisser (1959) correction for violation of the assumption of sphericity. The effects of Laterality or Distribution are reported when they interact with the experimental manipulations. Interactions between factors were followed up with simple-effect tests.

¹⁸ Indeed, as shown in Table 1, the size of the behavioral priming effect was minimal (e.g., 6 ms for 0° and 90° primes and 1 ms for 45° primes).

Figure 2

Schematic representation of the electrode montage and the channels included in statistical analyses.



Results

Behavioral Results

For the inferential analyses, we employed linear mixed-effects (LME) models in R (R Core Team, 2021) using the *lme4* (Bates et al., 2015, 2019) *lmerTest* (Kuznetsova et al., 2016) and *car* (Fox & Weisberg, 2019) packages. The models included two fixed factors: (1) prime-target relation (identity, unrelated) and (2) prime rotation angle (0°, 45°, and 90°). Error responses and very short RTs (less than 250 ms: 2 responses) were excluded from the RT analyses. The mean correct RTs and the accuracy in each condition are presented in Table 1. There were 11,762 observations. Because of the normality assumption required by LME analyses, the raw RTs were inverse-transformed ($-1000/\text{RT}$). For the accuracy analyses, the responses were coded as binary values (1 = correct, 0 = incorrect), and we used the *glmer* function in the *lme4* package. There were 12,480 observations. For both latency and accuracy analyses, we chose the maximal random effects model whose structure

converged¹⁹. The p -values for each main factor and the interactions were calculated using Wald χ^2 tests. In case of a significant interaction, we computed the effect of prime-target relation for each rotation angle with the *emmeans* package (Lenth et al., 2020).

Table 1

Mean correct response times (in ms) and accuracy (in brackets) for words

		Prime rotation		
		0°	45°	90°
Prime-target relation	Identity	578 (97.8)	597 (96.2)	605 (96.2)
	Unrelated	618 (96.9)	624 (96.4)	613 (96.4)
Priming effect (Unrelated – Identity)		40 (0.9)	27 (0.2)	8 (0.2)

Note: Priming effects for pseudowords were negligible (-6, 1, and 6 ms, for the 0°, 45°, and 90° primes, respectively).

Latency analyses. We found main effects of prime-target relation [$\chi^2(1) = 37.734$, $p < .001$] and angle [$\chi^2(1) = 26.345$, $p < .001$]. More importantly, the two factors interacted significantly [$\chi^2(2) = 31.537$, $p < .001$]. This interaction reflected that identity priming was significant for 0° primes (40 ms; 578 vs. 618 ms for identity and unrelated conditions, respectively; $t = -7.357$, $p < .001$) and 45° primes (27 ms; 597 vs. 624 ms for identity and unrelated conditions, respectively; $t = -4.977$, $p < .001$). In contrast, when the primes were rotated in 90°, there was only a non-significant 8-ms priming effect (8 ms; 605 vs. 613 ms, respectively; $|t| < 1$, $p = .64$).

Accuracy analyses. None of the factors or interactions approached significance (all $ps > .203$).

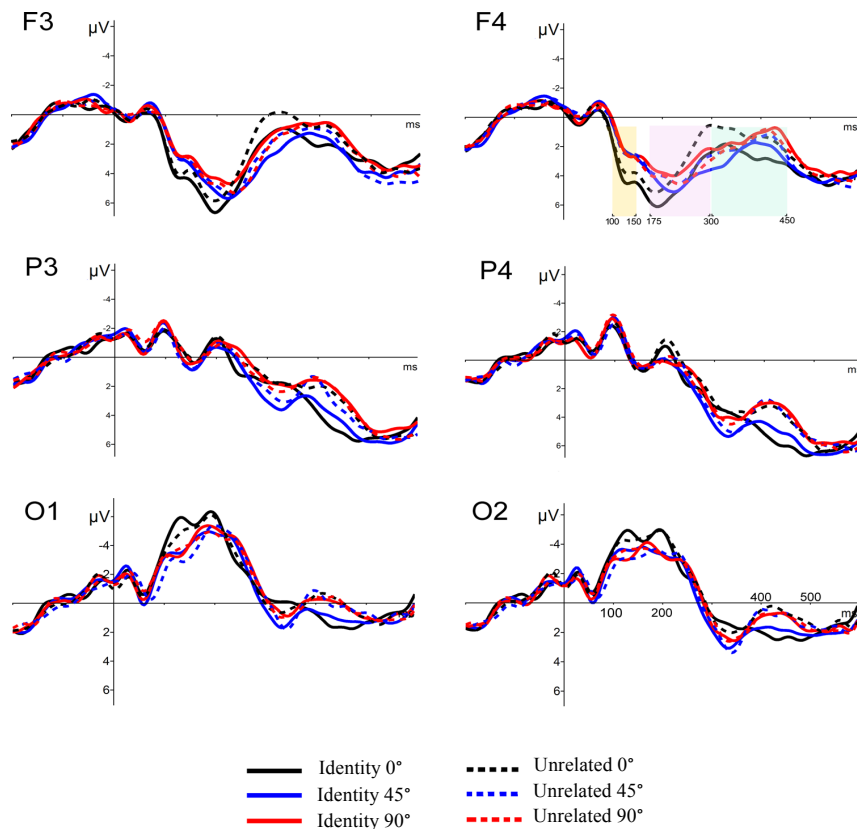
¹⁹ The maximal random effects models that converged were `lmer(-1000/RT ~ angle*relation + (1+relation|subject) + (1+relation|item))` and `glmer(accuracy ~ angle*relation + (1|subject) + (1|item))`

In sum, the latency data revealed that responses were faster for identity primes than for unrelated primes and, crucially, this difference was modulated by the rotation angle of the prime letters: it decreased as a function of rotations (0° [40 ms] > 45° [27 ms] > 90° [8 ms]).

ERP results

Figure 3

Grand average event-related potentials to words in the masked priming identity and unrelated conditions, for each word-letter rotation condition, in six representative electrodes of the areas of interest



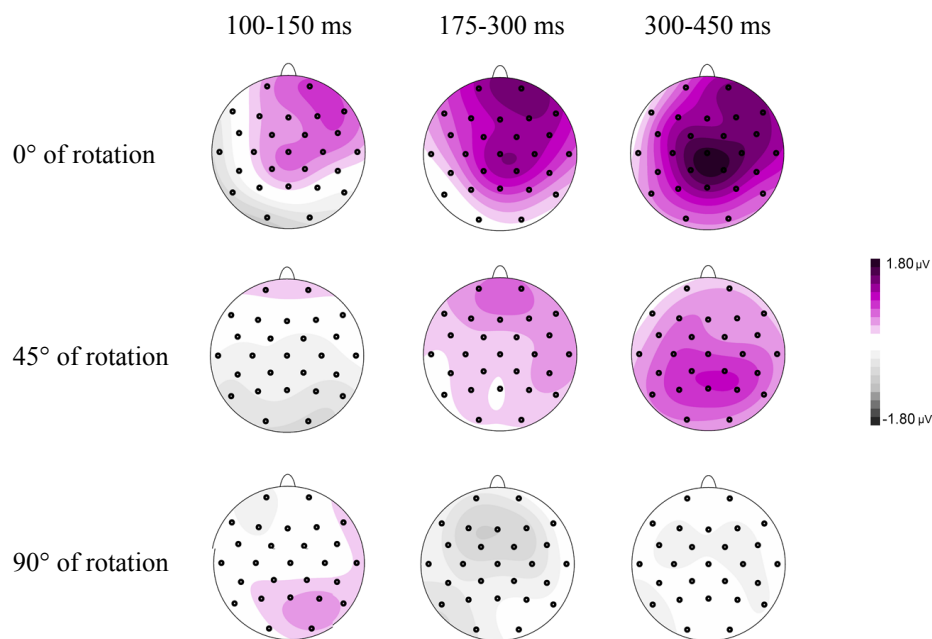
Note. Electrodes O1 and O2 show the negative counterpart of the N/P150. Color bars indicate the three time epochs under analysis.

Figure 3 shows the ERP waves for related and unrelated prime-target word pairs in six representative electrodes. Around 100 ms post-stimuli and over frontal electrodes, larger positive values are elicited by the 0° primes than for the 45° and 90° primes—this effect is inverted in the occipital electrodes (N/P150). For

the 0° primes, differences in amplitude between the identity and unrelated pairs start around 100 ms. This difference increases and peaks maximum of around 400 ms. For the condition with 45° rotation primes, the difference between identity and unrelated primes emerges later, around 200 ms of processing, and reaches its peak around 400 ms. Finally, for the 90° condition, there is no difference in amplitude between the identity and unrelated pairs. The evolution of the masked identity prime for each prime rotation is deployed in Figure 4. Below, we describe the results of the ANOVA for each time epoch.

Figure 4

Topographic distribution of the masked identity priming effect in the three time windows of the analysis.



Identity priming effect across letter rotation

Note. The priming effect was calculated as the difference in voltage amplitude between the event-related potential responses to identity vs. unrelated priming, for word targets preceded by each type of prime (0° rotation, 45° rotation, 90° rotation)

100- to 150 ms epoch. The ANOVA showed a main effect of rotation angle ($F[2, 40] = 4.37, p = .020$) that was modulated by the interaction with distribution ($F[4, 80] = 10.892, p < .001$). In frontal areas, larger positive values were observed for

the 0° primes compared to 45° primes ($A_{0^\circ} - A_{45^\circ} = 1.246$; $F[1, 20] = 19.94$, $p < .001$) and 90° primes ($A_{0^\circ} - A_{90^\circ} = 1.009$; $F[1, 20] = 9.76$, $p = .005$). In medial areas, we found the same pattern of data: larger positive values were observed for the 0° primes compared to 45° primes ($A_{0^\circ} - A_{45^\circ} = 0.902$; $F[1, 20] = 16.25$, $p = .001$) and 90° primes ($A_{0^\circ} - A_{90^\circ} = 0.725$; $F[1, 20] = 9$, $p = .007$); in both cases, there were no differences between 45° and 90° primes (both $F_s < 1$). The interaction between rotation, angle and relation reached significance in frontal and medial regions of the scalp ($F[4, 80] = 3.301$, $p = .050$): we found differences between identity and unrelated pairs that were restricted to 0° primes ($A_{ID0^\circ} - A_{UN0^\circ} = 1.045$ in anterior areas; $F[1, 20] = 8.06$, $p = .010$; $A_{ID0^\circ} - A_{UN0^\circ} = 0.752$ in medial areas ; $F[1, 20] = 5.17$, $p = .028$).

175- to 300 ms epoch. The ANOVA showed an effect of rotation angle that was modulated by laterality and distribution ($F[8, 160] = 4.62$, $p = .001$), a main effect of relation [$F(1, 20) = 4.90$, $p = .039$] and, more importantly, an interaction between rotation angle and relation ($F[2, 40] = 5.29$, $p = .011$). Regarding rotation, we found differences in amplitude between 0° and 45° primes in right-posterior regions ($F[1, 20] = 5.95$, $p = .024$) and between 45° and 90° primes in medial-anterior regions ($F[1, 20] = 5.49$, $p = .029$). Regarding the interaction between rotation angle and relation, for the 0° primes, we found larger positive values for identity than for unrelated pairs ($A_{ID0^\circ} - A_{UN0^\circ} = 1.142$; $F[1, 20] = 12.90$, $p = .002$). The difference between identity and unrelated pairs decreased for 45° primes ($A_{ID45^\circ} - A_{UN45^\circ} = 0.768$; $F[1, 20] = 5.03$, $p = .036$) and was absent for 90° primes ($A_{ID90^\circ} - A_{UN90^\circ} = -0.354$; $F < 1$).

300- to 450-ms epoch. The ANOVA showed a main effect of rotation angle ($F[2, 40] = 7.91$, $p = .002$), a main effect of relation ($F[1, 20] = 11.74$, $p = .003$), and an interaction between relation and rotation angle ($F[2, 40] = 5.337$, $p = .010$). Regarding rotation angle, follow-up analyses showed significant differences, restricted to identity pairs between 0° and 90° primes ($F[1, 20] = 21.61$, $p < .001$), and between 45° and 90° primes ($F[1, 20] = 26.69$, $p < .001$)—we did not find any differences between 0° and 45° primes ($F < 1$). Regarding relation, the largest difference in amplitude between identity and unrelated pairs was obtained for the

0° primes ($A_{ID0^\circ} - A_{UN0^\circ} = 1.345$; $F[1, 20] = 17.672$, $p < .001$). This difference was attenuated for the 45° primes ($A_{ID45^\circ} - A_{UN45^\circ} = 0.920$; $F[1, 20] = 7.343$, $p = .013$), and absent for 90° primes ($A_{ID90^\circ} - A_{UN90^\circ} = 0.202$; $F < 1$).

The time-course and size of the effects under study can be summarized as follows. First, the largest impact of rotation angle was observed in the N/P150, with the standard 0° prime orientation diverging from each two consecutive rotation angles. However, by 400 ms of target processing, brain responses for 90° rotated primes were the only ones that diverged from the canonical 0° and the 45° conditions. The effect of prime-target relation emerged by 150 ms for the standard prime orientation (0°). Importantly, the identity priming for 0° primes followed the natural course (i.e., it began around 100-150 ms and strengthened in the following processing stages). In contrast, the emergence of the priming effect for 45° primes was delayed, with its size also being reduced (i.e., it emerged around 170-300 ms and it was weaker than for the 0° primes). For 90° primes, the identity priming effect was absent in all components.

To sum up, the behavioral data showed that the masked identity priming effects were modulated by the rotation angle of the individual letters of the prime: it was strongest at 0° rotation angle (40 ms), decreased at the 45° rotation angle (27 ms), and vanished at the 90° rotation angle (8 ms). The ERP data showed an overall effect of rotation angle in the component associated with visual coarse feature encoding (N/P150): the amplitude was larger for the 0° primes than for the rotated primes (45° and 90°). Regarding masked identity priming, for the 0° primes, we already found effects in the N/P150 that substantially increased in the N250 and N400. For the 45° primes, the priming effect emerged later, in the N250, and increased in the N400—its size was weaker than for the 0° primes. Finally, the 90° primes did not produce identity priming in any ERP components. Thus, the present experiment revealed consistent and complementary findings on how letter-rotation affects masked identity priming in behavioral and electrophysiological measures.

Discussion

In the present ERP masked priming study, we examined the resilience of letter detectors to visual distortion via letter rotation. The LCD model (Dehaene et al., 2005), one of the leading neurally-inspired models of word recognition, posits that the letter detectors that convert the perceptual input into the abstract representation of letters would be severely hindered by letter rotations by around 40-45°. However, recent research using whole word rotations has questioned this claim. To shed light on this prediction with a more stringent scenario, we examined the temporal course of the effect of rotation of the words' constituent letters—instead of the whole words—in a masked identity priming paradigm (i.e., identity vs. unrelated primes; see Figure 1). The behavioral data showed that masked identity priming was modulated by letter rotation: it was strongest for the 0° primes (40 ms), decreased for the 45° primes (27 ms), and it was negligible for the 90° primes (8 ms).

More importantly, the ERP data revealed more fine-grained evidence of this pattern. First, in the N/P150, we found an overall effect of rotation consistency between prime and target, a result in line with previous findings (Vergara-Martínez et al., 2015; Gutierrez-Sigut et al., 2019; Petit et al., 2007)—note that neural activity at this stage of processing mainly deals with coarse physical features for the stimuli. Importantly, we also obtained a prime-target relation effect, but only when prime and target shared orientation (i.e., 0°). This pattern suggests that: (i) letter-form feature detectors are engaged very quickly during word processing, and (ii) the cost of letter rotation occurs very early, in a visual feature encoding stage.

In the N250 time window, we found differences between 0° and 45° primes, and 45° and 90° primes. More importantly, identity priming interacted significantly with prime rotation: The impact of primes on target processing was substantial for 0° primes. Moreover, identity priming effects were significant (although weaker) for 45° primes, whereas they were absent for 90° primes. Thus, at this point of processing, the ERP results were parallel to the behavioral data: identity priming was greatest for the 0° primes, reduced for 45° primes, and

null for 90° primes. Importantly, this pattern of ERPs data is consistent with the idea that, at this processing stage, letter size and shape invariance has already been achieved (Chauncey et al., 2008; Grainger & Holcomb, 2009). Nevertheless, the small identity priming effects for the 45° rotated primes (and the absence of priming for the 90° primes) in this time window offer evidence of the limited tolerance of letter detectors to rotation angle. In the first stages of visual-word recognition, the abstract orthographic representations of the primes are integrated with target processing, preferably if a coarse visual feature such as letter orientation is preserved. When the prime stimuli are presented in 0° (i.e., canonical orientation), the letter detectors of the following—identical—target are pre-activated, thus facilitating its processing. However, when the primes are presented in 45° or 90°, and the target in 0°, the pre-activation of the letter detectors decreases as a function of the orientation inconsistency between the prime and target letters. Thus, in this scenario, the processing of the identical target would be hindered.

Notably, the picture changes by 300-400 ms post-stimuli. We found no overall differences between 0° and 45° primes, but we found differences between 0° vs. 90° and 45° vs. 90°. More importantly, in the N400 window, larger identity priming occurred not only when the letters of the prime were presented at 0° but also—although to a lesser degree—when they were rotated 45°. This pattern suggests that the processing of the stimuli when the primes were presented at 0° followed the typical course (e.g., Vergara-Martínez et al., 2015). In contrast, the integration of orthographic information with 45° primes emerges later in processing. Thus, $\approx$ 45° rotated letters may need additional processing resources, such as top-down lexical mechanisms to help integrate the representations of primes and targets (see Benyhe & Csibri, 2020; Kim & Straková, 2012; Vergara-Martínez et al., 2021)²⁰.

²⁰ Parallel ANOVAs run on the ERPs for pseudoword targets revealed similar N/P150 effects of letter rotation to word targets ($F[2, 50] = 4.78, p = .018$; Distribution $\times$ Rotation: $F[2, 50] = 14.6317.83, p < .001$). Simple comparisons revealed significant differences between 0° and 45° over frontal ($F[2, 40] = 12.35, p = .002$) and central scalp areas ($F[2, 40] = 6.12, p = .022$). However, no effects of identity priming were observed in the time-epochs under study.

These findings are consistent with the findings reported by Kim and Straková (2012) in an unprimed lexical decision experiment. They obtained increased amplitudes in early time windows (P1 and N170) for 45° words relative to the canonical 0°, thus suggesting a cost in the initial moments of word-form analyses. Moreover, the present findings extend the results from the study conducted by Fernández-López et al. (2020) with parafoveal primes during sentence reading. They showed that skilled readers can efficiently convert parafoveal words composed of rotated letters to a stable orthographic code *only* when the rotation is $\leq 45^\circ$. Altogether, this pattern of results offers empirical support to the assumption of the LCD model (Dehaene et al., 2005) that letter detectors are hindered by letter rotation during the initial moments of processing for angles around 45°. Importantly, although it is beyond the scope of the present study, further work should directly examine the apparent discrepancies between the effects from the rotations of words as a whole (e.g., see Benyhe & Csibri, 2020) and the rotations of the individual letters within words.

In sum, the present masked priming study is the first to assess the electrophysiological brain signature of distortion in one single parameter, namely, letter orientation, during the early stages of word processing. Our results suggest that, during visual word recognition, participants can integrate the information from the primes composed of rotated letters up to 45°. However, this comes with a cost: the magnitude of the masked identity priming effects with 45° primes was noticeably smaller than with 0° primes in behavioral and ERP measures. Moreover, the fact that identity priming for 45° primes was weak in the N250 component suggests that rotations hampered the initial letter encoding and access to abstract representations (i.e., letter rotation impeded the automatic integration of the orthographic representation of the masked primes with the targets). Hence, our findings strongly suggest that the word processing system cannot easily handle letter rotations at around 45° or above, thus providing evidence supporting the predictions of the LCD model of word recognition (Dehaene et al., 2005).

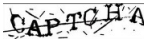
Letter rotations: Through the Magnifying Glass and What Evidence Found There

Abstract

Expert readers have a wide tolerance for distortions of the letters that make up a word. Nevertheless, the limits of this invariance are still under debate. To scrutinize this issue, we focused on a single parameter, letter rotation, as it serves to disentangle the predictions from neural-inspired models of word recognition. Whereas the Local-Combination-Detector (LCD) model predicts invariance up to 45° , the SERIOL model predicts a linear cost until 60° . To test these predictions, Experiments 1 and 2 employed four rotation angles (0° , 22.5° , 45° , 67.5°) in lexical decision and semantic categorization. The cost was minimal at 22.5° , sizeable at 45° , and considerably large at 67.5° . In Experiment 3, we focused on four moderate rotation angles ($<45^\circ$). We found a gradual reading cost that increased at 45° . Thus, while there is a resilience limit around 45° favouring LCD, less steep angles also produce a reading cost, backing the SERIOL model.

Key words: letter rotation; word frequency; lexical decision; semantic categorization

Introduction

One of the most remarkable skills of expert readers is their ability to quickly identify words despite their multiple variations in physical appearance. Indeed, heavily distorted words are used to differentiate between human and computer-based character extraction (e.g., ; see Hannagan et al., 2012, for evidence). Thus, it is not surprising that all leading models of visual-word recognition in alphabetical languages assume that the variations from the perceptual signal (e.g., *rabbit*, *rabbīt*, *rabbit*, Rabbit, RABBIT) are distilled into an abstract representation during the earliest moments of processing. In these models, the elements responsible for accessing lexical memory are layers of abstract letter detectors (e.g., *rabbit* and *rabbīt* would activate the same abstract letter units; see Grainger, 2018, for review).

The main goal of this paper is to examine to what degree the letter detectors in the word recognition system tolerate the degradation of the sensory input during lexical access. Given that the distortion of the physical signal of written text occurs along many different dimensions (i.e., all forms of *bad* penmanship can be bad in their own way), in this article, we focused on a single parameter of distortion: letter rotation. This decision has an added advantage: several neurally-inspired models of visual-word recognition (e.g., Local Combination Detector [LCD] model; Dehaene et al., 2005; Serial Encoding Regulated by Inputs to Oscillations within Letter units [SERIOL] model, Whitney, 2001) make precise claims on the impact of rotation angle on the resilience of letter detectors during word processing. In the following paragraphs, we offer a short overview of these two models, then we review the existing literature, and finally, we present a rationale for the three experiments reported in the paper.

The LCD model, postulated by Dehaene et al. (2005), is a hierarchical neurally-inspired model that explains word recognition via a series of local combinations of detectors that are sensitive to increasingly larger fragments of words, from shape fragments to complete words (see Figure 1 in Dehaene et al.,

2005). Crucially, at the level in which abstract representations of letters are recognized, the detectors are invariant to changes in size or format (i.e., a A a **A** and a share the same abstract units). Nevertheless, the model postulated that letter rotation hinders the normal mapping of the letter features: “letter detectors should be disrupted by rotation ($> 40^\circ$)” (Dehaene et al., 2005, p. 340). Specifically, the LCD model assumes that the visual word recognition system moves from the parallel encoding of letters to an effortful serial processing strategy for rotation angles above 40° - 45° . Notably, it is assumed that rotation angles below that boundary would induce “no measurable cost” (Vinckier et al., 2006, p. 2006; see also Cohen et al., 2008).

The SERIOL model (Whitney, 2001) is a theoretical framework that describes how word recognition is achieved via a series of processing layers. Specifically, the process of abstract letter identification depends on the activation of the nodes representing each letter, and this activation depends on the input provided by the letter feature level. The SERIOL model assumes that when the letters are rotated, the sensory input activates less the letter nodes, thus producing a reading cost as a function of rotation angles (i.e., degradation increases and the input decreases). Importantly, whereas this cost would be relatively small for moderate rotation angles, Whitney (2002) assumed that there would be a threshold at around 60° of rotation above which the input levels would not be sufficient to encode the word constituent letters automatically. Specifically, rotation angles above 60° would require extra top-down attentional processing (e.g., via activation in parietal areas in the brain) to guide a more letter-by-letter encoding. In this latter scenario, there would be a substantial delay in lexical access.

Although the LCD and SERIOL models share several principles concerning the role of rotation angle during word recognition, they vary in two crucial aspects. First, the boundary for the resilience of abstract letter detectors is lower in the LCD model than in the SERIOL model (around 45° vs. 60° , respectively). Second, in the LCD model, abstract letter detectors are assumed to be unaffected by moderate rotation angles (less than 40° - 45° ; an “invariance-with-limits perspective”, Kim & Straková, 2012). In contrast, the SERIOL model

predicts a gradual, increased cost as a function of rotation angle up to the boundary level of 60°.

In the present experiments, we aimed to disentangle the predictions of these models. Importantly, our interest is in the rotation of individual letters within words rather than the rotation of the whole word. The reason is that when a word is rotated as a whole, readers could rotate the entire stimulus to the canonical position and then process it as usual (i.e., one movement of mental rotation; see Whitney, 2002; Gómez & Perea, 2014, for discussion). Indeed, recent research has repeatedly shown that the rotation of the whole word produces sizable masked identity priming even at angles of 90° or greater (see Benyhe & Csibri, 2021; Perea et al., 2018; Yang & Lupker, 2019). By contrast, the rotation of the individual letters within a word involves the disruption of trans-letter features (i.e., features that are larger than letters but smaller than words; Mayall & Humphreys, 1996). Thus, the effect letter-by-letter word rotation represents a more stringent test of the LCD and SERIOL models than whole-word rotation.

Somewhat surprisingly, only a few studies have directly examined the claims of the LCD and SERIOL models concerning the resilience of letter detectors to the rotation of individual letters. One of the exceptions is the study conducted by Kim and Straková (2012). They recorded the event-related potentials (ERP) during a lexical decision task (i.e., a word vs. nonword discrimination task). The items' letters were presented in 0°, 22.5°, 45°, 67.5°, or 90°, and each word/nonword was presented for 200 ms. Kim and Straková focused on two early time windows (P1 [95-125 ms] and N170 [160-225 ms]). When considering the peak latencies, they found that angle rotation produced a linear increase on both P1 and N170 components. When considering amplitudes, they found a quadratic component of angle rotation for the P1 component: a rise between 0° and 22.5°, which remained stable until 67.5°, and then a decrease in amplitude at the 90° rotation. For the N170 component, they considered sub-windows. In the 160–190 ms interval, they found both an increase in linear and quadratic components from 0°-22.5° and 22.5°-45°, stable from 45°-67.5°, and a decrease from 67.5°-90°. In the 195-225 ms interval, they found a linear increase in all consecutive rotation angles.

Kim and Straková (2013) concluded that the increased amplitudes and delayed peaks with relatively moderate rotation angles (i.e., well below 45°) posed problems for the LCD model. However, a potential interpretive issue from the Kim and Straková (2013) experiment is that, to reduce the motor artifacts in the ERP signal, lexical decision responses had to be made after a cue presented 850 ms after the disappearance of the target. Therefore, they did not record a direct measure of the impact of letter rotation on the speed of lexical access; in other words, we cannot know whether the rotation effects are encapsulated in the encoding process, or instead, they cascade into the lexical processes. Indeed, they found remarkably similar word response times for 0° and 67.5° rotations. Thus, the effect of letter rotations in the ERPs at moderate angles (e.g., 22.5°) might be accounted for by changes in the visual appearance rather than a delay in lexical access (see Perea et al., 2015, for discussion).

In a more recent study, Blythe et al. (2019) measured participants' eye movements when reading sentences in which the individual letters within words were rotated (30°, 60°) or not. They found an overall reading cost on sentence reading for the words with 30° letters and an even larger one for the words with 60° letters. They also examined whether rotation angle interacted with word frequency by embedding a high- vs. low-frequency target word in each sentence. Their logic was that if rotation angle affected lexical processing, one would expect an interaction (i.e., greater word-frequency effects for the more degraded, 60° words). Results showed additive effects of rotation angle and word-frequency in the first-fixation duration on the target word (30° < 60°; high-frequency < low-frequency; i.e., the magnitude of word frequency effects was relatively consistent across all manipulations on these duration measures). Instead, gaze durations (i.e., the sum of fixation durations before leaving the target word) and total fixation durations (i.e., the sum of all fixations in the target word) showed an interaction between the two factors. Compared to the canonical 0° condition, the difference between high- and low-frequency target words increased in the 30° condition and even more in the 60° condition.

As reading was impaired even by relatively small rotations (i.e., 30°), Blythe et al. (2019) concluded that rotations hamper letter detectors at angles

below 45° during sentence reading and, as a result, the LCD predictions concerning letter rotations should be revised. Indeed, the gradual cost of letter rotation fits better with the SERIOL model. Nevertheless, the experiment conducted by Blythe et al. (2019) does not inform us as to whether the effect of letter rotation occurs in the first moments of processing (i.e., as posited by the LCD and SERIOL models) or at a later point during the integration processes—note that participants could have adjusted their eye movement pattern to the rotation angle of the sentences.

To shed light on this latter issue, Fernández-López et al. (2021) conducted a parafoveal preview experiment using Rayner's (1975) gaze-contingent boundary change paradigm during sentence reading. Crucially, the letters of the fixated target word and the rest of the sentence were in the upright orientation. The parafoveal preview was either identical or unrelated to the target word, and its letters were rotated 15°, 30°, 45°, or 60°. Results showed that the advantage of the identity preview condition in eye fixation times on the target word decreased progressively as a function of the rotation angle: the identity advantage was sizeable for 15° and 30°, weak for 45°, and absent for 60° (e.g., the benefit in single fixation durations was 26 ms, 13 ms, 7 ms, and 3 ms, respectively). Thus, the cost of letter rotation of parafoveal previews during sentence reading increased as a function of rotation angle, being substantial after 45°, providing empirical support to the LCD model. However, Fernández-López et al. (2021) measured the effects of the preview-on-target integration rather than directly measuring the impact of the rotation angle on eye movement measures. Furthermore, the previews were presented in the parafovea, where the quality of the spatial information is lower than in the fovea. As a result, it is difficult to completely ascertain whether the cost of letter rotations at 45° or above reflects the system bounds on early orthographic processing in a standard foveal scenario or whether it reflects structural limitations of parafoveal processing.

Taken together, previous research has shown that although there is a boundary beyond which processing words with rotated letters becomes effortful, relatively small rotation angles may also produce some reading cost. Although some of these findings favor the SERIOL model (e.g., a gradual cost of the

transpositions at rotation angles less than 40-45°), the evidence is scarce and subject to alternative interpretations (e.g., a modulation in early ERPs does not imply a cost in lexical access; see Vergara-Martínez et al., 2015, for discussion). Furthermore, the evidence on the existence of a specific boundary that critically hampers the workings of letter detectors and consequent word recognition is not conclusive either. Crucially, a limiting issue in most of the previous experiments on this issue was the lack of a lexical factor in their design—the exception was Blythe et al. (2019). As a result, it is complicated to ascertain whether the impact of letter rotation is perceptual (visually) or linguistically mediated. To infer the effect of letter rotation on lexical processing, we did include word frequency as a factor in the experiments.

In the present study, we conducted three experiments that examined in detail the cost of one type of perceptual degradation (letter rotation) on lexical access. In Experiment 1, we conducted a lexical decision experiment with high- and low-frequency words using rotation angles of 0°, 22.5°, 45°, and 67.5°. The manipulation of word-frequency allowed us to examine whether the impact of letter rotation is on an early encoding phase (i.e., additive effects of rotation angle and word-frequency), or whether its impact also carries on to subsequent lexical processing (i.e., more frequent words would be less affected by steeper rotations because of top-down lexical feedback). Experiment 2 was parallel to Experiment 1, except that we used a task that directly taps onto the access to lexical-semantic memory: semantic categorization (i.e., asking participants whether the presented item referred to an animal name or not). Notably, the use of different tasks in Experiment 1 (lexical decision) and Experiment 2 (semantic categorization) allowed us to explore whether the effect of angle rotation varies across tasks that tap into different processes. Finally, Experiment 3 was designed to examine in detail the effect of rotation angle in the region of 45° and below (i.e., 22.5°, 30°, 37.5°, and 45°), as this is the threshold boundary proposed by the LCD model.

We can envision three scenarios concerning Experiment 1 (lexical decision task). First, if there is a gradual, purely linear increment in the word response times across the various rotation angles (0°, 22.5°, 45°, and 67.5°), it would pose problems for both LCD and SERIOL models. Second, if word

response times are similar for the 0° and 22.5° rotation angles and then there is a dramatic increase at the 45° and 67.5° rotation angles, this outcome would favor the predictions of the LCD model. Third, a linear increase in word recognition times at the 0°, 22.5°, and 45° angles, and then a dramatic increase at the 67.5° angle, would favor the predictions of the SERIOL model. In addition, an interaction between rotation angle and word frequency would suggest that the effect of rotation angle spilled over lexical processes. In this latter scenario, top-down processes from the lexical level from high-frequency words may outweigh the harmful effects of perceptual degradation (see Vergara-Martínez et al., 2021, for evidence with handwritten words).

Experiment 1: Lexical Decision Task

Method

Participants

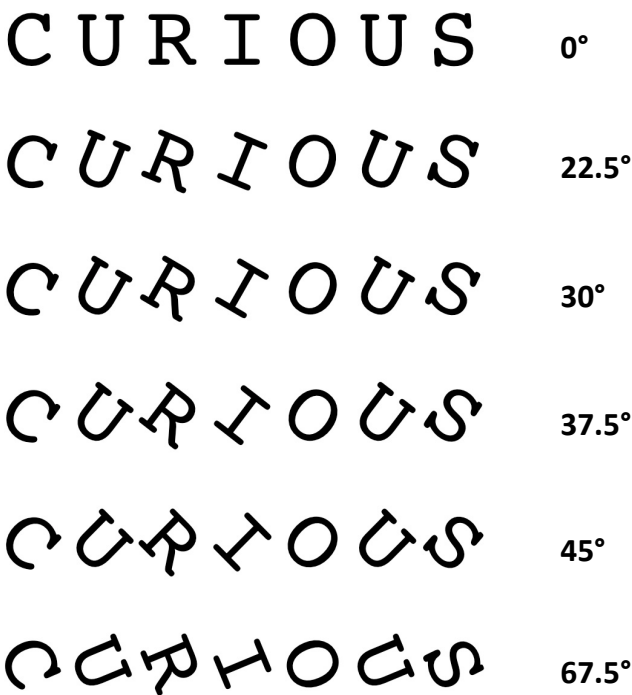
Sixty individuals took part in the experiment (mean age = 29 years-old; SD = 5.57; 39 women). With this sample size, we had 1,800 observations per condition, which is in line with the suggestion from Brysbaert and Stevens (2018) for small effects. Participants were recruited with Prolific Academic, a UK-based online crowd working platform (<http://prolific.ac>). Only English native speakers with no reading problems and normal/corrected vision could participate. All participants gave informed consent before the experiment and received monetary compensation. This and the following experiments obtained ethical approval from the Research Ethics Committee of the University of Valencia and followed the requirements of the Helsinki Convention.

Materials

We employed the 120 high-frequency and 120 low-frequency English words used by Aschenbrenner and Yap (2019). These two sets of words had been matched on a number of relevant psycholinguistic variables (e.g., number of letters, number of syllables, number of morphemes, OLD20, number of orthographic

neighbors, number of phonological neighbors, and PLD20), while differing (see Table 1 of Aschenbrenner & Yap, 2019) in SUBTLEX word frequency (Brysbaert & New, 2009). We also employed the 240 orthographically legal pseudowords selected by Aschenbrenner and Yap (2019)—these stimuli had been created with Wuggy (Keuleers & Brysbaert, 2010). The letters of each item (either words or pseudoword) could be presented in four different rotation angles (see Figure 1): 0° (i.e., canonical presentation), 22.5°, 45°, or 67.5°.

Figure 1
Rotation angles employed in the experiments.



Note. Letters in Experiment 1 and 2 were rotated 0°, 22.5°, 45°, and 67.5°. Letters in Experiment 3 were rotated 22.5°, 30°, 37.5°, and 45°.

Procedure

The script was written in Psychopy 3 software (Peirce & MacAskill, 2018) and was conducted online using Pavlovia (www.pavlovia.org). Before the experiment, all participants filled out a questionnaire with demographic data (age, gender, education level) via LimeSurvey (www.limesurvey.org). Participants were advised to do the experiment in a quiet room without distractions. Their task was

to decide whether the presented item was a word or not by pressing the “yes” (M) and “no” (Z) keys on the keyboard as quickly and accurately as possible. Within a given trial, a fixation cross was presented in the center of the screen for 500ms. Then, the target item was presented until a response was made or 2 sec had passed. We included 16 practice trials before the 480 experimental trials. There were brief breaks after every 100 trials. Altogether, the experiment took 18-20 minutes to complete.

Results and Discussion

For the analyses of the response times (RTs), we excluded the error responses (4.4% for words; 7.2% for pseudowords) and the very brief responses (less than 250 ms; 1 observation [less than 0.004%]). Time-outs after 2000 ms (0.17% for words, 0.55% for pseudowords) were coded as errors. Table 1 presents the mean correct RTs and error percentage in each experimental condition.

Table 1

Mean Response Times (in ms) and percent errors (in %) in each of the conditions in Experiment 1 (Lexical Decision Task)

		Rotation angle			
		0°	22.5°	45°	67.5°
Words	High Freq	608 (1.7)	608 (1.7)	657 (3.7)	781 (7.6)
	Low Freq	638 (4.0)	647 (3.4)	696 (4.7)	830 (8.7)
Frequency effect		30 (2.3)	39 (1.7)	39 (1)	49 (1.1)
Nonwords		719 (5.6)	748 (6.9)	819 (6.6)	985 (9.7)

The latency and the accuracy data were fitted with Bayesian linear-mixed effects models, using the *brms* package (Bürkner, 2017) in R (R Core Team, 2021). For the fits, we employed the exgaussian distribution for the latency data, thus capturing the skew of the RT data, and the Bernoulli distribution for the

binary accuracy data. We conducted separate analyses for word and nonword trials. The fixed effect factors for the word trials were Rotation angle of the items' letters (0°, 22.5°, 45°, and 67.5°, where the canonical 0° was the reference level) and Word-frequency (high vs. low; encoded as -0.5 and 0.5). For the nonword trials, the only fixed factor was Angle, but the general procedure was the same as that for the word trials. To examine the nature of the function of the reading cost due to rotation angle (i.e., linear, quadratic, cubic) in the latency data, we also conducted an analysis in which Rotation Angle was coded using polynomial contrasts. In all statistical analyses, we employed the maximal random-effect structure allowed in the design:

$$DV \sim \text{Rotation Angle} * \text{Word Frequency} + (1 + \text{Rotation Angle} * \text{Word Frequency} \mid \text{subject}) + (1 + \text{Rotation Angle} \mid \text{item}).$$

We carried out 5,000 iterations for each model—1,000 were for warmup—using four chains. All models fit adequately ($\hat{R} = 1.00$ in all cases). For the fixed effects, we report an estimate of the mean of the posterior distribution for each parameter (b parameter), together with its standard error (SE) and 95% credible interval (95% CrI) based on the quantiles of the distribution of the posterior estimates. We considered evidence for an effect when the 95% credible interval of its estimate did not cross zero.

Word stimuli

Response Times. We found evidence of an effect of word-frequency, $b = 24.13$, $SE = 4.08$, 95%CrI [16.03, 32.09]: responses were 39 ms faster for high-frequency words than for low-frequency words. Regarding the effect of rotation angle, words with 22.5° letters were responded to only minimally slower (5 ms) than those with 0° letters, $b = 4.95$, $SE = 2.85$, 95%CrI [-0.56, 10.48]. Words with 45° letters were responded 54 ms slower than those with 0°, $b = 41.32$, $SE = 3.55$, 95%CrI [34.35, 48.30]. Finally, words with 67.5° letters were responded to 183 ms slower than those with 0°, $b = 116.80$, $SE = 6.51$, 95%CrI [103.84, 129.60]. None of the rotation angle effects interacted with word-frequency (all $bs < 7.71$; all 95%CrIs crossed zero).

The trend analyses showed clear evidence of linear and quadratic components of rotation angle (linear: $b = 86.26$, $SE = 4.81$, 95%CrI [76.95, 95.81], quadratic: $b = 35.25$, $SE = 3.05$, 95%CrI [29.33, 41.27]). These effects did not interact with word-frequency (all b values of the interaction were less than 5 ms and their 95%CrIs crossed zero).

Accuracy. We found a main effect of word-frequency, $b = -0.94$, $SE = 0.27$, 95%CrI [-1.49, -0.40]: responses were 2% less error prone for high-frequency than for low-frequency words. Concerning the effect of rotation angle, we did not find differences between the accuracy on responses to 0° and 22.5° words, $b = 0.34$, $SE = 0.37$, 95%CrI [-0.36, 1.13]—this effect did not interact with word-frequency, $b = 0.15$, $SE = 0.36$, 95%CrI [-0.55, 0.84]. Moreover, words with 45° letters were responded a 7% less accurately than those with 0° , $b = -0.74$, $SE = 0.29$, 95%CrI [-1.31, -0.16]—this effect interacted with word-frequency, $b = 0.74$, $SE = 0.32$, 95%CrI [0.11, 1.38], showing that the word-frequency effect was smaller with 45° letters. Finally, words composed of 67.5° letters were responded a 4.7% less accurately than those with 0° , $b = -1.67$, $SE = 0.27$, 95%CrI [-2.20, -1.16]—this effect also interacted with word-frequency, $b = 0.78$, $SE = 0.31$, 95%CrI [0.18, 1.39], showing that the word-frequency effect was smaller with 67.5° letters.

Nonword stimuli

Response Time analysis. Pseudowords composed of 22.5° letters were responded 29 ms more slowly than those with 0° , $b = 16.33$, $SE = 2.81$, 95%CrI [10.84, 21.88]. The latency cost increased to 100 ms for the pseudowords with 45° letters, $b = 56.38$, $SE = 4.26$, 95%CrI [48.04, 64.74], and even more, 266 ms, for the pseudowords with 67.5° letters, $b = 143.49$, $SE = 9.85$, 95%CrI [123.89, 163.11].

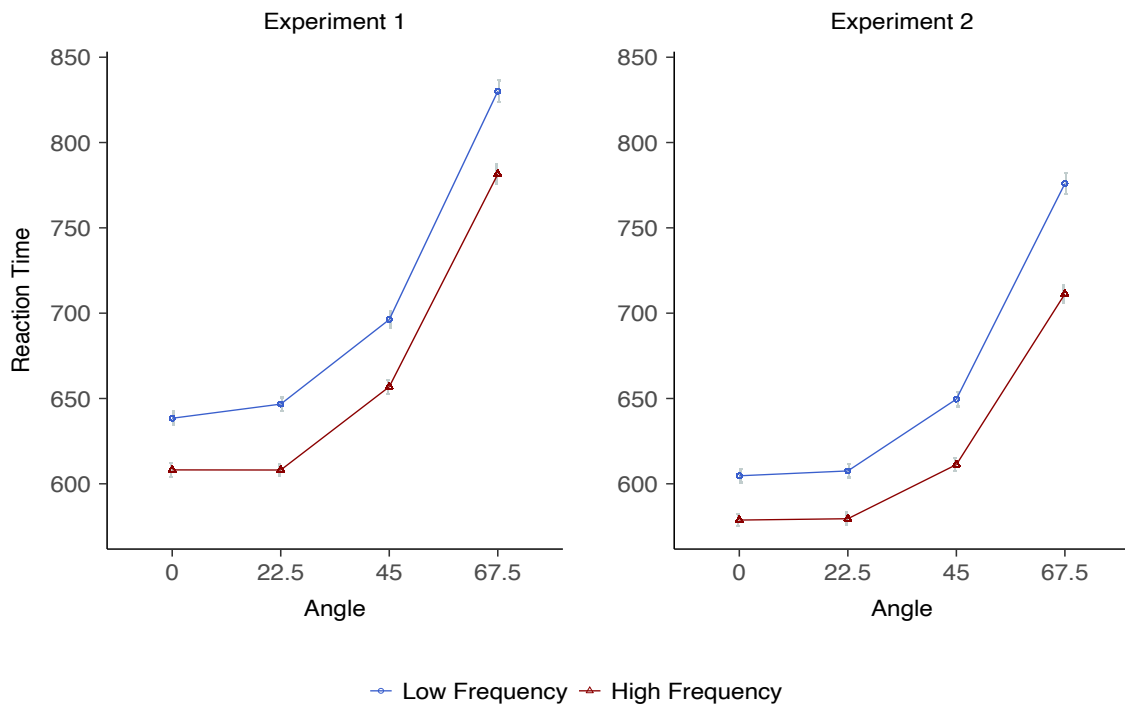
The trend analyses showed clear evidence of linear and quadratic components of rotation angle (linear: $b = 105.30$, $SE = 6.97$, 95%CrI [91.77, 119.02], quadratic: $b = 35.31$, $SE = 3.77$, 95%CrI [27.85, 42.66]). We also found some evidence of a small cubic component, $b = 5.38$, $SE = 2.34$, 95%CrI [0.83, 9.90].

Accuracy analysis. We did not find evidence of differences in accuracies between 0° and 22.5°, $b = -0.17$, $SE = 0.15$, 95%CrI [-0.46, 0.12], and 0° and 45°, $b = -0.21$, $SE = 0.15$, 95%CrI [-0.49, 0.09]. Nevertheless, responses to pseudowords with 67.5° letters were 3.1% more error prone than responses to pseudowords with 0°, $b = -0.63$, $SE = 0.14$, 95%CrI [-0.90, -0.34].

The present experiment revealed that word identification times increased nonlinearly at the steeper rotation angles: when compared to the canonical 0° format, the reading cost was minimal for the words with 22.5° letters (4 ms), it increased to 54 ms for the words with 45° letters, and considerably more, to 183 ms, for the words with 67.5° letters. This reading cost was not accompanied by a substantially larger word-frequency effect (e.g., the word-frequency went from 30 ms at the canonical 0° format to 49 ms at the 67.5° angle) (see Gomez & Perea, 2014, for a similar pattern when rotating whole words in a lexical decision task). Overall, these increases can be modeled by a combination of linear and quadratic components (see Figure 2). The analyses of the error rates showed a similar pattern: the percentage of errors was similar for the 0° (2.9%) and 22.5° (2.6%) angles, and then increased for the 45° angle (4.2%) and even more for the 67.5° angle (8.2%).

Figure 2

Response times (in ms) for high- and low-frequency words across the different rotation angles in Experiment 1 (Lexical Decision Task) and Experiment 2 (Semantic Categorization)



The pseudoword data revealed an analogous pattern, except that the effects of rotation angle were steeper than those for words. We found a sizeable cost (29 ms) for the pseudowords with 22.5° letters compared to those with 0° letters. Relative to this canonical 0° format, the cost increased to a much larger degree for the 45° (100 ms) and the 67.5° (266 ms). Thus, rotation angle hinders pseudowords to a greater degree than words. To obtain statistical support of this dissociation, we conducted an analysis with Lexicality (word vs. pseudoword) and Rotation angle as polynomial contrasts. We found evidence of a greater cost for pseudowords than for words in the linear component (interaction Angle x Lexicality) $b = 12.53$, $SE = 6.22$, 95%CrI [0.38, 24.71]).

In sum, the present experiment showed a sizable reading cost of rotation angle for words at the 45° angle (54 ms; 1.4% errors) and, much more

dramatically, at the 67.5° angle (183 ms; 5.3% errors), but not at the 22.5° angle (5 ms; -0.3% errors). This pattern in word responses favors the predictions of the LCD model. However, pseudowords showed greater reading costs than words, including the 22.5° rotation angle (29 ms; 1.3% errors). One explanation for this dissociation is that top-down feedback from the lexical level could have partly overridden the cost of letter rotation at moderate angles.

The question now is whether rotation angle also affects a recognition task that requires word identification at a semantic level. Experiment 2 was designed to examine the resilience of word recognition to letter rotation using a task that is less sensitive to purely visual elements and more on access to meaning in lexical memory (semantic categorization task) (see Perea et al., 2020, for discussion).

Experiment 2: Semantic Categorization Task

Method

Participants

The number of participants, sampled from the same population as Experiment 1, was increased to 80 (mean age = 29 years old; SD = 5.41; 60 women) so that the number of observations per condition was comparable to that of Experiment 1—note that the number of trials per condition was smaller than in Experiment 1 (see the Materials subsection below).

Materials

For the non-animal words, we used a subset of 184 words from the set of high- and low-frequency words from Experiment 1 (92 high-frequency words and 92 low-frequency words). The reason is that we excluded those words which referred to animal names or were closely related to animals (e.g., mother,

mountain, food). We also selected a set of 92 words for the animal words, thus keeping the 2:1 ratio used in previous experiments (e.g., Perea et al., 2020).

Procedure

It was parallel to Experiment 1, except that the participants were asked to decide whether the item referred to an animal name (press “M”) or not (press “Z”).

Results and Discussion

Table 2

Mean Response Times (in ms) and percent errors (in %) in each of the conditions in Experiment 2 (Semantic Categorization Task)

		Rotation angle			
		0°	22.5°	45°	67.5°
Non-Animal	High Freq	579 (1.6)	580 (1.0)	611 (1.4)	711 (2.7)
	Low Freq	605 (1.8)	608 (1.8)	650 (2.6)	776 (4.0)
Frequency effect		26 (0.2)	28 (0.8)	39 (1.2)	65 (1.3)
Animal		618 (9.7)	629 (11.1)	652 (9.4)	709 (11.0)

The analyses plan was the same as in Experiment 1. Although the main focus was on non-animal words (i.e., the set of high vs. low-frequency words from Experiment 1), we also analyzed the data from the animal words for completeness. For the non-animal words, the fixed factors were Angle and Word-frequency; for the animal words, the fixed factor was Angle. Table 2 displays the mean RT and error rate in each experimental condition.

Non-animal words

Response Time analysis. We found that responses were 40 ms faster for high- than for low-frequency words., $b = 19.44$, $SE = 3.86$, 95%CrI [11.94, 26.96]: Regarding rotation, words with 22.5° letters were responded only 2 ms slower

than those with 0°, $b = 2.18$, $SE = 2.98$, 95%CrI [-3.71, 8.07]—this effect did not interact with word-frequency, $b = -0.17$, $SE = 4.25$, 95%CrI [-8.46, 8.11]. Words with 45° letters were responded 39 ms more slowly than those with 0°, $b = 41.32$, $SE = 3.54$, 95%CrI [17.52, 31.40]—this effect did not interact with word-frequency, $b = 8.03$, $SE = 4.88$, 95%CrI [-1.52, 17.56]. Finally, words with 67.5° letters were responded to 152 ms slower than those with 0°, $b = 85.19$, $SE = 6.61$, 95%CrI [72.07, 98.24]—this effect did interact with word frequency, $b = 24.83$, $SE = 8.40$, 95%CrI [8.22, 41.27]: the difference between responses to high and low frequency words was 39 ms greater when the letters were rotated 67.5° than when presented at 0°.

The trend analyses showed evidence of both linear and quadratic components of rotation angle (linear: $b = 61.71$, $SE = 4.68$, 95%CrI [52.43, 70.89], quadratic: linear: $b = 29.15$, $SE = 3.22$, 95%CrI [22.85, 35.46]). In addition, word-frequency interacted with the linear component of rotation angle, $b = 17.93$, $SE = 5.68$, 95%CrI [6.80, 29.02]—the evidence of a quadratic component was weak (i.e., 95%CrI crossed zero), $b = 8.02$, $SE = 4.39$, 95%CrI [-0.50, 16.74].

Accuracy analysis. We did not find evidence of an effect of word-frequency, $b = -0.06$, $SE = 0.33$, 95%CrI [-0.70, 0.59]. Regarding rotation, words with 22.5° letters and words with 45° letters, were responded as accurately as those with 0°, (0° vs 22.5°, $b = 0.71$, $SE = 0.40$, 95%CrI [-0.02, 1.55]; 0° vs 45° $b = 0.34$, $SE = 0.36$, 95%CrI [-0.33, 1.10]). However, responses to words with 67.5° letters were 2.2% more error prone than responses to words with 0°, $b = -0.61$, $SE = 0.30$, 95%CrI [-1.21, -0.02]. None of the rotation angle effects interacted with word-frequency (all $bs < 0.54$; all 95%CrIs crossed zero).

Animal words

Response Time analysis. Relative to the words with 0° letters, we found a cost that increased with rotation angle: 11 ms for 22.5°, $b = 11.48$, $SE = 2.65$, 95%CrI [6.31, 16.66]; 34 ms for 45°, $b = 27.41$, $SE = 3.09$, 95%CrI [21.36, 33.48]; and 91 ms for 67.5°, $b = 64.76$, $SE = 4.11$, 95%CrI [56.75, 72.84].

The trend analyses revealed evidence of both linear and quadratic components of rotation angle (linear: $b = 46.87$, $SE = 3.22$, 95%CrI [40.70, 53.27], quadratic: $b = 12.78$, $SE = 2.12$, 95%CrI [8.59, 17.00]).

Accuracy analysis. We did not find evidence of differences in accuracy relative to the words with 0° in any of the three angles (all $bs < 0.23$; all 95%CrIs crossed zero).

The pattern of word recognition times for the non-animal words in the semantic categorization task mimicked that of the word recognition times in the lexical decision task. There was a minimal cost from 0° to 22.5° rotation angles (2 ms) that increased at the 45° rotation angle (39 ms) and was even more dramatic at the 67.5° rotation angle (152 ms). The effect of word-frequency only increased at the 67.5° angle relative to the canonical 0° format (65 vs. 26 ms, respectively). Error rates on non-animal names were very low (2.1%) and only showed a small disadvantage in the 67.5° angle relative to the 0° angle (3.4% vs. 1.7%, respectively).

The pattern of data for animal words was similar to that for non-animal words except that there was a small reading cost for the words composed of 22.5° letters relative to the canonical format (11 ms); this cost increased for the 45° angle (34 ms) and even more for the 67.5° angle (91 ms). The error rates did not show differences across rotation angles.

In sum, the effect of rotation angle in the range between 0° and 67.5° is remarkably similar in lexical decision and semantic categorization tasks, reflecting both linear and quadratic components. The reading cost from 0° to 22.5° rotation angle was minimal, sizeable for the words with 45° letters, and substantial for the 67.5° angle (see Figure 2 for a comparison of the two experiments).

While the overall pattern in Experiments 1 and 2 is generally consistent with the claims of Dehaene et al.'s LCD model, we found some gradual increases as a result of rotation angle in some conditions, even for moderate rotations that deserve further scrutiny. To further study this issue, we designed Experiment 3. The goal was to zoom in on the effects of letter rotation in angles up to 45° in a

semantic categorization task. For this purpose, in Experiment 3, we chose four angles of rotation whose difference was almost indiscernible (i.e., 7.5°): 22.5°, 30°, 37.5°, and 45°.

The key question in Experiment 3 was whether there is an “abrupt” boundary at around 45° in which rotation angle makes word recognition substantially less automatic, or whether there is a gradual linear cost as a function of rotation angle at moderate angles. If there is a critical boundary level at around 40-45° at which letter detectors cannot efficiently deal the sensory input—as predicted by the LCD model, we would expect negligible increases from the 22.5° up to 37.5° rotation angle and then a large increase at the 45° angle (i.e., a quadratic component). Alternatively, rotation angles below 45° could produce a small but steady accumulative reading cost in each step—as predicted by the SERIOL model. We would expect a linear (but not quadratic) component of the rotation angle in this latter case.

Experiment 3: Semantic Categorization Task

Method

Participants

We recruited 80 new participants (mean age = 29 years old; SD = 5.42; 66 women) from the same population as in the previous experiments—note that the sample size was the same as in Experiment 2.

Materials

They were the same as in Experiment 2, except that the rotation angles were 22.5°, 30°, 37.5°, and 45°.

Procedure

It was the same as in Experiment 2.

Results and Discussion

The analysis plan was the same as in Experiment 2. Table 3 presents the averages per condition in the latency and error data.

Table 3

Mean Response Times (in ms) and percent errors (in %) in each of the conditions in Experiment 3 (Semantic Categorization Task)

		Rotation angle			
		22.5°	30°	37.5°	45°
Non-Animal	High Freq	588 (0.9)	583 (1.3)	590 (0.7)	606 (1.2)
	Low Freq	600 (1.4)	615 (1.1)	614 (1.5)	639 (2.0)
Frequency effect		12 (0.5)	32 (-0.2)	24 (0.8)	33 (0.8)
Animal		620 (10.0)	625 (9.1)	628 (9.3)	643 (8.6)

Non-animal words

Response Time analysis. We found evidence of an overall effect of word-frequency, $b = 17.13$, $SE = 3.91$, 95%CrI [9.56, 24.90]: responses to low-frequency words were 25 ms slower than responses to high-frequency words. Regarding the rotation angle effect, when comparing to responses with 22.5° rotated letters, words with 30° letters produced a minimal cost (5 ms), $b = 2.26$, $SE = 2.98$, 95%CrI [-3.13, 7.58], that slightly grew for words with 37.5° rotated letters (8 ms), $b = 8.41$, $SE = 2.81$, 95%CrI [2.95, 13.98]. Importantly, for words with 45° rotated letters, the cost increased sharply (29 ms), $b = 20.08$, $SE = 3.08$, 95%CrI [14.04, 26.21]. None of these effects interacted with word-frequency (all $bs < 4$, all 95%CrI crossed zero).

This pattern was confirmed by the trend analyses, which showed evidence of linear and quadratic components of rotation angle (linear: $b = 14.75$, $SE = 2.27$,

95%CrI [10.29, 19.22], quadratic: $b = 4.72$, $SE = 2.01$, 95%CrI [0.76, 8.67]). There were no signs of an interaction of rotation angle with word-frequency (all b 's < 2.27).

Accuracy analysis. We did not find evidence of an effect of word-frequency, $b = -0.38$, $SE = 0.41$, 95%CrI [-1.19, 0.42]. We did not find evidence of an effect of rotation relative to 22.5° letters in any of the three angles (all b s < 0.39, all 95%CrI crossed zero) or an interaction between word-frequency and rotation angle (all b s < 0.83, all 95%CrI crossed zero).

Animal words

Response Time analysis. Relative to the words with 22.5° letters, there were no differences when the letters were rotated at 30° (5 ms), $b = 2.55$, $SE = 2.65$, 95%CrI [-1.38, 8.63]. We found some cost of rotation at 37.5° (8 ms), $b = 6.97$, $SE = 2.77$, 95%CrI [1.54, 12.41], which abruptly increased at the angle of 45° (23 ms), $b = 18.60$, $SE = 2.56$, 95%CrI [13.62, 23.65].

The trend analyses again revealed evidence of both linear and quadratic components (linear: $b = 13.16$, $SE = 2.09$, 95%CrI [9.01, 17.34], quadratic: linear: $b = 4.18$, $SE = 2.10$, 95%CrI [0.11, 8.33]), but not a cubic component ($b = .91$, $SE = 1.78$, 95%CrI [-1.59, 5.42]).

Accuracy analysis. There were no differences in accuracy of the various angle conditions relative to the 22.5° words (all b s < 0.30, all 95%CrI crossed zero).

The present experiment showed small increases in the word response times across rotation angles, except for the more abrupt change at the 45° rotation angles. Compared to the 22.5° baseline, the recognition times for non-animal words increased 5, 8, and 29 ms when the letters were rotated 30°, 37.5°, and 45°. We found a parallel pattern for animal words: recognition times increased in 5, 8, and 23 ms. This pattern reveals a combination of linear and quadratic components, suggesting that (1) there is something special at a rotation angle of around 45° (i.e., the quadratic component predicted by the LCD model) and (2) there is a gradual cost in lower rotation angles (i.e., the linear component posited by the SERIOL model).

General Discussion

In the present study, we designed three experiments to test the predictions of the LCD and SERIOL models concerning the resilience of letter detectors to rotation at relatively moderate angles. The LCD model (Dehaene et al., 2005) assumes invariance up to around rotation angles around 40°-45° after which the resilience of letter detectors drops abruptly, whereas the SERIOL model (Whitney, 2002) postulates a gradual increase of word processing cost with rotation angle (at least until an abrupt shift at around 60°).

When using a relatively large range of rotation angles (0°, 22.5°, 45°, and 67.5°; i.e., a difference of 22.5° in each step), we found a linear cost accompanied by a quadratic cost in the word response times in lexical decision (Experiment 1) and semantic categorization (Experiment 2). Relative to the canonical format (0°), the cost was minimal for the 22.5° words, substantial for the 45° words, and dramatically large for the 67.5° words (see Figure 2). Notably, in the lexical decision task (Experiment 1), we found that the effect of letter rotation was larger for pseudowords (0° vs. 22.5°: 29 ms; 22.5° vs. 45°: 100 ms; 45° vs. 67.5°: 266 ms) than for words (0° vs. 22.5°: 5 ms; 22.5° vs. 45°: 54 ms; 45° vs. 67.5°: 183 ms). This pattern suggests that lexicality helps to cope with degraded stimuli. To further examine the complexities of rotation angles close to the critical threshold proposed by the LCD model (i.e., 45°), we conducted a third experiment using rotation angles at or below that threshold (i.e., 22.5°, 30°, 37.5°, and 45°; a difference of 7.5° in each step) with a semantic categorization task. Results showed both a gradual linear increase—favoring the SERIOL model, and an extra bump for the 45°—as predicted by the LCD model.

Taken together, the present findings are consistent with different aspects of both the LCD and SERIOL models. There is a gradual cost of rotation angle that increases at 45°, and this cost is dramatically amplified at 67.5° (see Figure 2). Thus, on the one hand, as predicted by the LCD model, there is an abrupt cost in word recognition times when the letters are rotated at 45° relative to more moderate values (i.e., 22.5°, 30°, or 37.5°). On the other hand, as predicted by the SERIOL model, an angle of rotation of 45° is not an all-or-none edge, and

moderate angles also generate a cost. Of note, we also found a dramatic increase in word recognition times in Experiments 1-2 at the 67.5° angle. Indeed, the word-frequency effect in the semantic categorization task of Experiment 2 was substantially greater at a 67.5° angle relative to the canonical format (65 vs. 26 ms, respectively). In contrast, the word-frequency effect only slightly increased for moderate angles (i.e., the effect was 28 and 38 ms at the 30° and 45° angles, respectively)—the pattern was similar (but weaker) in the lexical decision task (Experiment 1). Thus, at a steep angle, like 67.5°, participants might have engaged in extra top-down processing to help access lexical-semantic information—as predicted by both LCD and SERIOL models. Thus, the claims of these models may complement rather than negate each other.

The present data fit well with previous research across various paradigms. In a lexical decision experiment, Kim and Straková (2013) found that increasing letter-rotation led to monotonic increases in the P1 and N170 (i.e., ERPs that mark the beginning of word-form analysis) amplitudes up to 45°-67.5°. Notably, they found an effect of letter rotation relative to the canonical format for 22.5° words (i.e., well below 45°), and suggested that moderate deviation from preferred stimulus properties (<45°-67.5°) could recruit additional processing resources to cope with distorted stimuli. Likewise, during sentence reading, Blythe et al. (2019) found a reading cost in fixation durations when letters were rotated 30° relative to the canonical 0° format—this effect was dramatically larger when for 60° words. Similarly, in a parafoveal preview paradigm during sentence reading, Fernández-López et al. (2021) found that the benefit of parafoveal previews was sizeable at rotation angles of 15°. This parafoveal benefit decreased at 30°, even more so at 45°, and was absent at 60°.

Thus, considering the previous and present empirical evidence across a variety of paradigms on the processing of rotated letters, we propose that (i) even moderate letter rotations produce a (small) cost in word processing; (ii) this cost increases nonlinearly, while being manageable, at around 45°, and (iii) this cost increases dramatically around 60° or larger. In the latter context, individuals are likely to use more serial processing of stimuli, through increased parietal

activation, to cope with the visual distortion caused by the rotation angle (see Cohen et al., 2008; Whitney, 2010).

To sum up, we conducted three experiments to shed light on the processing of letter rotation and discerning between the predictions of two neurally-inspired models of word recognition (LCD model, SERIOL model). We found that the cost elicited by rotation angle increases smoothly at angles below 45° , as proposed by the SERIOL model. The letter detectors reach a resilience limit at around 45° (consistent with the LCD model), and the cost becomes dramatic at steeper angles (e.g., 67.5°). We also found that lexical factors could partly outweigh the harmful effects of letter rotation. We propose combining the assumptions from models to build a more comprehensive explanation of how rotation angle, and more generally, visual distortion, hampers word recognition.

Conclusions

The literate brain contains specialized mechanisms that are exquisitely tuned to recognize written words. An example is that we can read words and texts distorted in different dimensions. The present section included the studies that aimed to test the tolerance of the word recognition system to letter rotations with three different methodological approaches. As indicated in the Introduction, the added value of examining the impact of letter rotations is that leading neural models of word recognitions have explicit predictions about how the system copes with them. While the Local Combination Detectors (LCD) model (Dehaene et al., 2005) predicts invariance up to 45°, the SERIOL model (Whitney, 2001, 2002) predicts a linear cost until a threshold around 60°.

In the first study, *How resilient is reading to letter rotations? A parafoveal preview investigation* (Fernández-López et al., 2021), we aimed to examine the predictions of the LCD and SERIOL models in a normal reading scenario. We conducted a reading experiment with sentences in uppercase using Rayner's (1975) gaze-contingent boundary change paradigm in which we examined whether identity parafoveal previews were more effective than unrelated parafoveal previews when the previews were rotated at 15°, 30°, 45°, or 60°. Importantly, apart from the parafoveal previews, all text, including target words, was presented with letters in their canonical upright orientation. We found a remarkably similar pattern of findings for all the fixation measures that we assessed on the target word (i.e., first-fixation duration, single fixation duration, and gaze duration). The advantage of the identity condition was greatest when the preview letters were rotated 15°, and it was numerically smaller but still robust, at 30°. For the 45° rotations, the advantage of the identity preview condition was only marginal. Finally, for the 60° rotations, there were no signs of a benefit of the identity preview condition. This pattern of results suggested that, although letter detectors appear to be severely disrupted at rotation angles of 45° or larger in the very early moments of processing, as posited by the LCD model,

there was also a gradual cost for letter rotations below that boundary, as postulated by the SERIOL model.

In the second study, *On the time course of the resilience of letter detectors to rotations: A masked priming ERP investigation*, we examined the temporal course of the effect of rotation of the words' constituent letters in a masked identity priming paradigm. We focused on testing the prediction of the LCD model: "letter detectors should be disrupted by rotation ($> 40^\circ$)" (Dehaene et al., 2005, p. 340). Specifically, the target was presented in the canonical orientation, and the primes' letters were presented in angles of 0° , 45° , or 90° . The behavioral data showed that masked identity priming was modulated by letter rotation: it was strongest for the 0° primes, decreased for the 45° primes, and it was negligible for the 90° primes. The Event-Related Potentials (ERP) data allowed us to examine in detail the time course of processing of words with rotated letters. Amplitude comparisons showed that the effect of prime-target relation emerged by 150 ms for the standard prime orientation (0°). Importantly, the identity priming for 0° primes followed the natural course (i.e., it began around 100-150 ms and strengthened in the following processing stages). In contrast, the emergence of the priming effect for 45° primes was delayed, with its size also being reduced (i.e., it emerged around 170-300 ms and it was weaker than for the 0° primes). For 90° primes, the identity priming effect was absent in all components. This pattern strongly suggested that letter rotations at around 45° impeded the automatic integration of the orthographic representation of the masked primes with the targets, thus providing evidence supporting the predictions of the LCD model of word recognition.

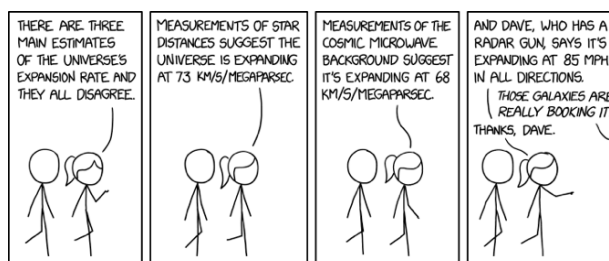
In the third study, *Letter rotations: Through the Magnifying Glass and What Evidence Found There*, we conducted three experiments that scrutinized the cost of letter rotation on lexical access and tested the predictions of the LCD and SERIOL models concerning the resilience of letter detectors to relatively moderate angles of letter rotation. In Experiment 1, we conducted a lexical decision task with letter stimuli rotated 0° , 22.5° , 45° , and 67.5° . In Experiments 2 and 3, we conducted a semantic categorization task with letters rotated 0° ,

22.5°, 45°, and 67.5° (Experiment 2) and 22.5°, 30°, 37.5°, 45° (Experiment 3). Results in Experiments 1 and 2 showed that, relative to the canonical format (0°), the cost was minimal for the 22.5° words, substantial for the 45° words, and dramatically large for the 67.5° words. Importantly, in the lexical decision task (Experiment 1), we found that the effect of letter rotation was larger for pseudowords than for words. Moreover, in Experiment 2, we found that letter rotation interacted with word frequency but only for words with 67.5° words—a similar (but numerically smaller) trend occurred in the lexical decision task. This pattern suggested that individuals might have recruited top-down lexical processes to cope with rotations. Finally, Experiment 3 found a gradual reading cost that increased at the 45° angle. Taken together, these results show that the cost elicited by rotation angle increases smoothly at angles below 45°, as proposed by the SERIOL model. The letter detectors reach a resilience limit at around 45° (consistent with the LCD model), and the cost becomes dramatic at steeper angles (e.g., 67.5°).

Thus, the three studies of this section aimed to test the tolerance of the word recognition system to visual distortion via the examination of letter rotation. The added value of examining letter rotations is that two leading models of word recognition have different predictions about how the brain copes with them. Considering the outcomes of the three studies and building up a comprehensive picture, the main ideas that can be drawn are summarized as follows: (i) the cost of letter rotation originates very early, in the first stages of word recognition. Importantly, (ii) even moderate rotations lead to a cost. Nevertheless, (iii) this cost is easily overcome, and the system can handle rotations up to 45° in the initial moments of processing. Importantly, (iv) to cope with steeper rotations, individuals seem to engage extra top-down (lexical) feedback. In sum (v), we found that the cost elicited by rotation angle increases smoothly at angles below 45°, as proposed by the SERIOL model. The letter detectors reach a resilience limit at around 45° (consistent with the LCD model), and the cost becomes dramatic at steeper angles. Therefore, (vi) the most suitable option is combining the assumptions from SERIOL and LCD models to build a more comprehensive

explanation of how rotation angle, and more generally, visual distortion, hampers word recognition.

Methodological contributions to the visual word recognition research



© XKDC

The more significant breakthrough in studying the reading process started around the late 1950s, coinciding with what is often called the cognitive revolution (see Pollatsek & Treiman, 2015, for review). This revolution broke with the behaviorist approach and led to many ingenious new techniques for studying reading and word identification processes, including the lexical decision task (Rubinstein et al., 1970, 1971a, 1971b) and the masked priming procedure (Forster & Davis, 1984). In the lexical decision task, participants are asked to distinguish as fast as possible words from nonwords. The time to determine that a letter string like `dog` is a real word provides researchers an index of how long it takes to contact the neural representation of that word in the reader's internal dictionary. If we want to tap on the early stages of word processing, we should use the priming technique, considered "the major methodological innovation of the last two decades" (Grainger, 2008). In this technique, a mask is presented for 500 ms, then replaced by a prime for 50 ms, so participants cannot see it, which is replaced by the target stimulus. Importantly, researchers usually manipulate the prime-target relationship, which can be related or unrelated. Participants' responses are often assessed with behavioral measures, like reaction time and accuracy, but also with more direct measures of brain activity, such as that obtained with functional magnetic resonance imaging (fMRI) and electroencephalogram (EEG).

Notably, the combination of the lexical decision and the masked priming has led to much empirical work investigating letter and word processing. Nevertheless, the technique requires further examination to build firm conclusions on the outcomes obtained with masked priming. Here, I present two studies that attempted to understand the cognitive underpinnings of masked priming and one of its variants, the sandwich priming. In the first study, entitled *Can response congruency effects be obtained in masked priming lexical decision?* (Fernández-López et al., 2019), we aimed to test whether participants apply the same instructions to the primes and targets in the masked priming lexical decision task by examining the response congruency effect with unrelated primes. The response congruency effect in lexical decision can be defined as the difference in response times between a congruent trial (i.e., when the prime elicits the same response as the target; e.g., `bright-WINDOW` or `geuzel-POLCIR`) and an incongruent trial (e.g., `geuzel-WINDOW` or `bright-POLCIR`). Importantly, the examination of this effect allows us to scrutinize “the state of the lexicon after the prime stimulus has been processed up to an early stage” (Loth & Davis, 2010b, p. 174). Specifically, we examined whether there were differences in the influence of a masked unrelated wordlike vs. nonwordlike primes (e.g., `geuzu` vs. `xdfpk`) on the processing of the target stimuli. An added value of this comparison is that it allowed us to test the predictions of leading models of word identification, the Bayesian Reader model (Norris & Kinoshita, 2008) and the family of interactive activation models (i.e., the Interactive Activation model from McClelland & Rumelhart [1981] and its successors [Multiple Read-Out Model, Grainger & Jacobs, 1996; Dual Route Cascade Model, Coltheart et al., 2001; Spatial Coding Model, Davis, 2010]).

Despite the numerous advantages of the masked priming technique, it is not sensitive enough to detect subtle manipulations (e.g., `lght-LIGHT` vs. `liht-LIGHT`; Grainger, 2008) or examine the letter and word recognition processes in special populations, such as developing readers or readers with dyslexia. The reason is that, because of the short stimulus-onset asynchrony between prime and target, the magnitude of masked priming effects is small, in the range of 10-50 ms in behavioral studies. To overcome this limitation, Lupker

and Davis introduced the sandwich technique in 2009. This technique is a variant of the standard masked priming in which the addition of a pre-prime, identical to the target, magnifies the size of priming effects (see Campos et al., 2020; Colombo, et al., 2020; Gutiérrez-Sigut et al., 2017; Ktori et al., 2014; Lupker et al., 2020a, 2020b; Lupker et al., 2015; Meade et al., 2020; Ziegler et al., 2014; Taikh & Lupker, 2020). Nevertheless, the specific mechanisms at play and how they differ from the standard technique, have not been fully elucidated (Lupker & Davis, 2009 vs. Trifonova & Adelman, 2018). In the second study of this section, *Unveiling the boost in the sandwich priming technique* (Fernández-López et al., 2021), we examined whether the priming effects in the sandwich technique are simply magnified or, instead, the underlying processes are fundamentally altered.

Can response congruency effects be obtained in masked priming lexical decision?

Abstract

In the past decades, researchers have conducted a myriad of masked priming lexical decision experiments aimed at unveiling the early processes underlying lexical access. A relatively overlooked question is whether a masked unrelated wordlike/unwordlike prime influences the processing of the target stimuli. If participants apply to the primes the same instructions as to the targets, one would predict a response congruency effect (e.g., *book*-TRUE faster than *fiok*-TRUE). Critically, the Bayesian Reader model predicts that there should be no effects of response congruency in masked priming lexical decision, whereas interactive-activation models offer more flexible predictions. We conducted three masked priming lexical decision experiments with four unrelated priming conditions differing in lexical status and wordlikeness (high-frequency word; low-frequency word; orthographically legal pseudoword; consonant string). Experiment 1 used wordlike nonwords as foils, Experiment 2 used illegal nonwords as foils, and Experiment 3 used orthographically legal hermit nonwords as foils. When the foils were orthographically legal (Experiments 1 and 3; i.e., a standard lexical decision scenario), lexical decision responses were not affected by the lexical status or wordlikeness of the unrelated primes, as predicted by the Bayesian Reader model and the selective inhibition hypothesis in interactive-activation models. When the foils were illegal (Experiment 2), consonant-string primes produced the slowest responses for word targets and the fastest responses for nonword targets. The Bayesian Reader model can capture this pattern assuming that participants in Experiment 2 were making an orthographic legality decision (i.e., anything legal must be a word) rather than a lexical decision.

Key words: lexical decision; masked priming; word recognition

Introduction

Since its inception in 1984, Forster and Davis' masked priming technique has become a fundamental tool to examine the initial moments of letter and word recognition in cognitive psychology (see Forster, 2013; Grainger, 2008) and cognitive neuroscience (e.g., ERPs: Grainger & Holcomb, 2009; MEG: Lehtonen et al., 2011; fMRI: Dehaene et al., 2001). In the standard setup, a briefly presented lowercase prime is preceded by a pattern mask and followed by an uppercase target stimulus until the participant's response, so that readers are not usually aware of the prime. This allows researchers to examine various types of prime-target relationships (e.g., identity priming: *steal*-STEAL vs. *horse*-STEAL; orthographic priming: *steam*-STEAL vs. *green*-STEAL; phonological priming: *steel*-STEAL vs. *shell*-STEAL; morphological priming: *stole*-STEAL vs. *brink*-STEAL; associative/semantic priming: *thief*-STEAL vs. *nurse*-STEAL; translation priming: *robar*-STEAL vs. *venir*-STEAL). Importantly, this technique minimizes the strategic processes that may occur with visible, unmasked primes (see Dehaene et al., 2001, for fMRI evidence; see Gomez et al., 2013, for modeling evidence).

Whilst the masked priming paradigm has been employed with a number of laboratory word identification tasks (e.g., lexical decision, naming, semantic categorization, same-different), most masked priming experiments have used the lexical decision task. Indeed, researchers have compiled large masked priming datasets with this task (see Adelman et al., 2014, for a mega-study of masked priming lexical decision using 27 form-related priming conditions) and most contemporary models of visual word recognition offer quantitative predictions at the behavioral level (i.e., response times [RTs] and accuracy) of masked priming effects in the lexical decision task (e.g., Dual-Route Cascaded model: Coltheart et al., 2001; Spatial Coding Model: Davis, 2010; Multiple Read-Out Model: Grainger & Jacobs, 1996; Bayesian Reader model: Norris & Kinoshita, 2008).

In the present paper, we focus on a relatively neglected phenomenon in masked priming lexical decision that may allow us to test the predictions of leading models of visual word recognition: the response congruency effect with

unrelated primes. The response congruency effect in lexical decision can be defined as the difference in response times between a congruent trial (i.e., when the prime elicits the same response as the target; e.g., *bright-WINDOW* or *geuzel-POLCIR*) and an incongruent trial (e.g., *geuzel-WINDOW* or *bright-POLCIR*). The idea is that participants may apply the task instructions (e.g., an implicit categorization of the “yes” vs. “no” decision) not just to the targets but also to the masked primes. As Loth and Davis (2010b) indicated, the examination of response congruency effects in masked priming lexical decision allows us to scrutinize “the state of the lexicon after the prime stimulus has been processed up to an early stage” (p. 174). Thus, this effect may allow us to elucidate the processes underlying masked priming lexical decision.

Response congruency effects have been reported in other masked priming tasks. For instance, Naccache and Dehaene (2001; see also Dehaene et al., 1998) found faster responses in a numerical categorization task (“is the number higher than 5?”) when the masked prime is from the same category as the target (e.g., prime: 7; target: 9) than when it is not (e.g., prime: 2; target: 9). Importantly, Dehaene et al. (1998) found evidence of a differential effect on motor preparation for congruent and incongruent trials using the lateralized readiness potential as a marker, thus favoring the idea that participants make implicit decisions to the prime stimuli. Similarly, Forster (2004) found a response congruency effect in a masked semantic categorization task for narrow categories (e.g., *answer-WINDOW* faster than *collie-WINDOW* for the category “type of dog”). To explain this response congruency effect, Forster (2004) stated: “an implicit decision about the category membership of the prime is reached before an explicit decision about the target is reached” (p. 283). If this implicit decision also takes place in masked priming lexical decision, one would expect faster “yes” responses to words after an unrelated word prime than after an unrelated nonword prime—and conversely, one would expect faster “no” responses to nonwords after an unrelated nonword prime than after an unrelated word prime. Critically, a leading computational model of visual word recognition, namely, the Bayesian Reader model (Norris, 2006; see also Norris & Kinoshita, 2008) makes a straight prediction: there should be no effects of response congruency in masked priming

lexical decision. The predictions from interactive activation models are less straightforward, although the presence/absence of a response congruency effect may be informative to decide whether there is homogeneous or selective inhibition at the lexical level.

The Bayesian Reader model (henceforth, BR model; Norris, 2006) assumes that words are represented in an orthographic/lexical perceptual space. Lexical decisions are based on the comparison of the evidence that the visual input is produced—or not—by a word. Information that increases the odds that the stimulus is a word will decrease the odds that the stimulus is a nonword until a criterion is achieved and a response is triggered. To simulate masked priming lexical decision, the BR model assumes that “evidence from both the prime and targets are integrated in reaching a decision” (Kinoshita & Norris, 2012, p. 3). The masked prime generates evidence that the target is in a particular area of the perceptual space, thus predicting faster lexical decision responses to orthographically related pairs (e.g., *trie*-TRUE) than to control pairs (*fiok*-TRUE). Critically, if prime and target do not share any orthographic features, the prime will not alter the evidence for the target in lexical decision—note that in Norris and Kinoshita’s account of masked priming lexical decision, the information from the prime is not sufficient to provide specific evidence that the target is a word or not. That is, neither the masked word prime *food* nor the masked pseudoword prime *fiok* would have an effect on the responses to the target word TRUE because they are far apart in perceptual space. As Norris and Kinoshita (2008) claimed, “there should be no response-congruence effects in lexical decision” (p. 442). What we should also note, however, is that the Bayesian Reader model assumes that priming effects depend on the decision required by the task, thus not precluding an effect of response congruence in other tasks. For instance, the Bayesian Reader model can capture an effect of response congruence in masked priming semantic categorization tasks (e.g., “is the word an animal name?”). In a semantic categorization task, “the evidence consists of distributed semantic features that are diagnostic of category membership” (de Wit & Kinoshita, p. 1738). Specifically, semantic features like <builds nests> in the masked prime *eagle* would constitute evidence toward a

“yes” decision, whereas semantic features like <made of metal> in the masked prime `stapler` would provide evidence toward a “no” decision, thus producing a response congruency effect.

Alternatively, the Interactive Activation (henceforth IA) model (McClelland & Rumelhart, 1981) and its successors (Multiple Read-Out Model, Grainger & Jacobs, 1996; Dual Route Cascade Model, Coltheart et al., 2001; Spatial Coding Model, Davis, 2010) assume that all units at the word level receive activation from the letter and visual feature levels. When the activation level of a word unit exceeds its resting level—which is a function of word-frequency—it sends inhibition to all other word units. In the IA model, a “yes” response in lexical decision occurs when the activation level of a word unit exceeds some asymptotic activation threshold (see Jacobs & Grainger, 1992, for the first implementation of the IA model to lexical decision), whereas a “no” response would be produced when none of the word units exceeds threshold after a certain processing duration (see Coltheart et al., 1977). The IA model can easily capture masked orthographic priming effects: the processing of a prime (e.g., the pseudoword `trie`) would partially activate similarly spelled word units (e.g., `true`, `tree`, and `trip`), thus predicting faster lexical decision times to `trie-TRUE` than to `fiok-TRUE` (see Perea & Rosa, 2000, for simulations). Using the original setting of the IA model, the activated word units inhibit all other word units, regardless of whether they are orthographically similar or not (i.e., homogeneous lateral inhibition). That is, upon presentation of the masked prime `food`, its corresponding word unit would send inhibitory signals to other word units (e.g. `TRUE`). This inhibition would occur to a lesser degree when the prime is not a word (e.g., the prime `fiok` would produce less inhibition on the word unit `TRUE` than the prime `food`). This specification of the IA model predicts slower “yes” lexical decision times for those target words preceded by an unrelated word prime than when preceded by an unrelated pseudoword prime (e.g., `food-TRUE` slower than `fiok-TRUE`). More generally, the model predicts that “yes” lexical decision times for those target words preceded by an unrelated prime should be a direct function of the lexical activation generated by the prime: high-frequency word prime (`food-TRUE`) > low-frequency word prime (`glop-TRUE`) >

pseudoword prime (*fiok*-TRUE) > consonant-string prime (*ghdk*-TRUE) (see Bailyss et al., 2009, for simulations showing this exact pattern). Thus, the original IA model predicts a response congruency effect—or rather an effect of the wordlikeness of the unrelated primes—for “yes” responses.

Importantly, the predictions in the IA model are dependent on the nature of the parameter settings. Davis and Lupker (2006) proposed a selective inhibition mechanism in which the activated word units would only inhibit those word units that are close in lexical space (i.e., similarly spelled word units; e.g., *tree*, *tire*, *tore*, *free* for the word unit *true*, but not for unrelated word units like *food*). In this scenario, neither the presentation of the high-frequency word prime *food* nor the presentation of the consonant-string prime *ghdk* would modify the activation level of the target word *TRUE*. Therefore, the selective inhibition hypothesis in the IA model would predict a null effect of response congruency for “yes” responses (see Bailyss et al., 2009, for simulations).

Finally, the predictions on “no” responses in lexical decision in the IA model may depend on how these responses are performed. If the temporal deadline for “no” responses is adjusted after an unrelated prime, as suggested by Forster (1998), one would not expect any response congruency effects for nonwords. That is, the temporal deadline for the pseudoword *UFEZ* would be the same regardless of whether the unrelated prime is a word (e.g., *true*) or not (e.g., *tybe*).

Table 1

Summary of masked priming lexical decision experiments with different types of unrelated primes

Research paper	Target type	Mean RT (in ms)					(*)
		Prime type					
		HFW	H/LFW	LFW	PW	NW	
Sereno (1991)	HF word	621			623		
	LF word			694	702		
	Pseudoword		689		678		
Perea et al. (1998)	HF word	676		681	679		
	Pseudoword	810		811	810		
Norris & Kinoshita (2008, Exp. 1)	HF word	561			561		
	LF word			634	621		
	Pseudoword		672		682		
Bailys et al. (2009, Exp. 1)	HF word	597		600	602	615	
	LF word	688		669	679	686	
Bailys et al. (2009, Exp. 2)	HF word	581		584			
	LF word	668		666			
Loth & Davis (2010a, Exp. 4)	MF word		636		638		
	Pseudoword		794		793		
Perea et al. (2010)	Word		635		641		
	Pseudoword		697		695		
Perea et al. (2014, Exp. 1)	HF/LF word		678		682		
	Pseudoword		727		726		
Jacobs et al. (1995, Exp. 2) (Items repeated several times)	HF word	493			513		*
	Pseudoword	515			498		*
Loth & Davis (2010a, Exp. 1)	HF word	475				498	*
	Illegal nonword	520				497	*
Loth & Davis (2010a, Exp. 2)	HF word	531				539	*
	Pseudoword	586				576	*
Loth & Davis (2010a, Exp. 3)	HF word	588				601	*
	Pseudoword	713				721	
Perea et al. (2014, Exp. 2) (Items repeated several times)	HF word	590			601		*
	Pseudoword	668			665		

Note. (1) * Refers to a reported effect from unrelated primes (2) HFW: high frequency word; H/LFW: high/low frequency word; LFW: low frequency word; MF: medium frequency; NW: illegal nonword; PW: pseudoword; (3) Wordlikeness of nonwords targets from Loth and Davis (2010a): Experiment 4 > Experiment 3 > Experiment 2.

The empirical evidence of an effect of response congruency—or more generally an effect of the wordlikeness of the unrelated prime—in masked priming lexical decision is scarce and non-conclusive (see Table 1 for a summary). While most experiments failed to obtain an effect of the unrelated primes (Baillyss et al., 2009; Loth & Davis, 2010a, Experiment 4; Norris & Kinoshita, 2008; Perea et al., 1998; 2010; 2014, Experiment 1; Sereno, 1991), several experiments did find an effect (e.g., Jacobs et al., 1995; Loth & Davis, 2010a, Experiments 1, 2 and 3; Perea et al., 2014, Experiment 2). There are two procedural differences between these two types of outcomes. The null findings occurred when the target stimuli were presented once and the foils were orthographically legal (i.e., the most usual scenario in word recognition experiments). In contrast, the experiments that reported an effect from the unrelated primes either repeated the target stimuli several times (Jacobs et al., 1995; Perea et al., 2014, Experiment 2) or the foils were orthographically illegal (e.g., SYKDD; Loth & Davis, 2010a). To re-examine in depth this issue, we focused on the response congruency effect when the target stimuli are presented only once (i.e., the most common setting in word recognition experiments) and we varied the wordlikeness of the nonword targets: orthographically legal nonwords matched in sublexical characteristics with the words in Experiment 1, orthographically illegal nonwords in Experiment 2, and orthographically legal nonwords with no close neighbors in Experiment 3.

The main goal of the present study was to examine in detail under which conditions the effect of response congruency could occur in masked priming lexical decision. Unlike previous published experiments, which focused exclusively on very few types of orthographically legal primes or on a single comparison of orthographically legal vs. orthographically illegal primes, we chose four types of unrelated primes that differed in wordlikeness (60 items per condition): 1) a high-frequency word (e.g., línea-JAULA, coche-GAULO [the English for line-CAGE, car-GAULO]); 2) a low-frequency word (e.g., censo-

JAULA, boina-GAULO [census-CAGE, beret-GAULO]); 3) an orthographically legal pseudoword (e.g., bunte-JAULA, vomon-GAULO [bunte-CAGE, vomon-GAULO]); and 4) a consonant string (e.g., pdtnr-JAULA, lbctr-GAULO [pdtnr-CAGE, lbctr-GAULO]). Thus, there were two unrelated word priming conditions and two unrelated nonword priming conditions. This allows us not only to examine the overall response congruency effect but also effect of the frequency of the word primes (high-frequency vs. low-frequency) and the effect of legality of the nonword primes (orthographically legal pseudowords vs. consonant strings).

Importantly, we also examined whether the difficulty of the word/nonword decision could modulate the response congruency effect in masked priming lexical decision. The rationale here is that the processes underlying lexical decision move faster toward the “yes” or “no” boundaries (i.e., evidence for one response or the other is accumulated faster) when words and nonwords are easily discriminable (e.g., using orthographically illegal nonwords as foils) (see Ratcliff et al., 2004, for empirical and modeling evidence). This may maximize the chances to observe an effect due to an implicit decision on the primes. For instance, the high-frequency prime *world* might be implicitly categorized as “yes”, and this would speed up the responses to the target word *CHAIR* relative to the consonant-string prime *wtshn*. However, this prime-induced congruency effect may be more difficult to detect when the nonword foils are orthographically legal (see Loth & Davis, 2010b for a similar reasoning). (One might argue that the size of effects is usually diminished when response times are faster, but masked priming effects tend to be similar in magnitude for fast and slow responses; see Gomez et al., 2013). In Experiment 1, the foils were orthographically legal nonwords matched in subsyllabic elements with the target words (e.g., *RERTA*). This is the typical scenario in word recognition experiments. In Experiment 2, the word/nonword discrimination was made substantially easier by having orthographically illegal nonwords (*DPTME*) as foils, and hence maximizing the chances of a prime-induced congruency effect on target stimuli from an implicit “yes” or “no” decision to the prime. To anticipate the results, we found no signs of an effect from the unrelated primes in Experiment 1, whereas

in Experiment 2 we found an effect of the unrelated primes on target performance, but only for consonant-string primes. Thus, one could argue that this latter finding may have been driven by a legality judgment decision (“if the target is orthographically legal, say yes”) rather than by lexical decision. To examine whether the effects found in Experiment 2 were due to change in the nature of the task (i.e., orthographic legality rather than lexical decision) or because the words and nonwords were easily discriminable, we designed a third experiment. In Experiment 3, the foils had no close neighbors (e.g., GAZUR) but they were always orthographically legal. Therefore, while the word/nonword discrimination in Experiment 3 is easier than in Experiment 1, lexical decisions need to be made in terms of lexical activation rather than orthographic legality.

According to the BR model (Norris & Kinoshita, 2008), one would expect no effects of the lexicality/frequency of the unrelated primes in masked priming lexical decision for word and nonword targets. Therefore, the presence of an effect from the unrelated primes would pose some problems for Norris and Kinoshita’s (2008) account of masked priming lexical decision—a similar null effect is predicted by Forster’ (2004) account of masked priming lexical decision, in which lexical decisions cannot be generated by masked primes. The predictions from the original IA model hinge on whether they assume selective inhibition or not; in former case, the model predicts no effects from the different types of unrelated primes, whereas in the latter, the model predicts a direct relation between the lexical decision times and the wordlikeness of the unrelated primes (see Bailyss et al., 2009, for simulations). We defer a thorough description of more sophisticated mechanisms for lexical decision responses in other IA models (i.e., MROM and SCM).

Experiment 1

Method

Participants

Twenty-four students from the Universitat de València, all of them native speakers of Spanish with normal/corrected-to-normal vision and with no history of reading disorders, took part voluntarily in the experiment. This study was approved by the Experimental Research Ethics Committee of the Universitat de València (Spain). Before starting the experiment, all participants signed an informed consent form.

Materials

We selected 240 five-letter target words from the Spanish subtitle database (Duchon et al., 2013). The mean frequency was 16 occurrences per million (range 5-40)—this corresponded to an average Zipf frequency of 4.15 (range 3.74-4.61) (see van Heuven et al., 2014, for the advantages of using Zipf frequencies when describing word information). The mean OLD20 (Yarkoni et al., 2008) was 1.41 (range: 1.00-2.65). We also created a set of 240 nonwords to act as foils. These nonwords were created from the target words using Wuggy (Keuleers & Brysbaert, 2010). To act as unrelated primes, we selected 120 high-frequency words ($M = 201$ occurrences per million, range 49-1304; mean Zipf frequency = 5.14, range 4.69-6.11) and 120 low-frequency words ($M = 2$ occurrences per million, range 0-5; mean Zipf frequency = 3.26, range 2.17-3.70)—all of them of five letters. Each target stimulus, presented in uppercase, was preceded by a lowercase prime that could be: 1) an unrelated high-frequency word; 2) an unrelated low-frequency word; 3) a pseudoword (obtained using Wuggy [50% from high-frequency words and 50% from low-frequency words]); or 4) a five-letter consonant string. To rotate the four priming conditions across all target words, we created four counterbalanced lists following a Latin square design. To familiarize the participants with the task, the 480-trial experimental phase (i.e., 60 items per condition) was preceded by a short practice phase composed of 16

trials (8 word trials and 8 nonword trials) with the same characteristics as the experimental trials—these practice trials were not included in the analyses.

Procedure

Each participant was tested alone in a quiet room. DMDX software (Forster & Forster, 2003) was used to display the sequence of stimuli and to register the timing/accuracy of the responses. In each trial, a pattern mask (i.e., a series of five #'s) was displayed for 500 ms in the center of a computer screen. The mask was replaced by a lowercase prime stimulus for 50 ms, and then the prime was replaced by an uppercase target stimulus. The target stimulus was displayed on the screen until the participant responded or 2 seconds had passed. All stimuli were presented in black (Courier New 14-pt font) on a white background. Participants were instructed to decide whether the target stimulus was a Spanish word or not by pressing the “yes” or the “no” key. Both speed and precision were stressed in the instructions. Instructions did not mention the existence of any lowercase stimuli. Sixteen practice trials preceded the 480 experimental trials. Participants did not receive feedback on reaction times or error rates during the experiment. The session lasted for around 18-22 minutes.

Results and discussion

Error responses and extremely short RTs (less than 250 ms: 4 responses) were omitted from the latency analyses—note that the deadline to respond was set to 2 sec. The mean RTs for the correct responses and the accuracy in each experimental condition are presented in Table 2.

To examine the influence of unrelated primes on the processing of word/nonword targets, we employed linear mixed effects (LME) models in R (R Core Team, 2017) using the *lme4* (Bates et al., 2015) and *lmerTest* packages (Kuznetsova et al., 2016) with two fixed factors: target lexicality (word, nonword) and prime type (high frequency word, low frequency word, pseudoword, consonant string). For type of prime, we created three orthogonal contrasts following the three research questions: (a) the effect of prime lexicality (word prime [high frequency word, low frequency word] vs. nonword [pseudoword, consonant string]); (b) the effect of prime word-frequency (high frequency word-prime vs. low frequency word-

prime); and (c) the effect of the wordlikeness of the nonword primes (pseudoword vs. consonant string). Because of the normality assumption required by LME analyses, the raw RTs were inverse-transformed ($-1000/\text{RT}$). There were 10,957 observations in the latency analyses. We chose the maximal random effects model whose structure converged. (Using untransformed RTs in the LMEs or using ANOVAs produced exactly the same pattern of results as that reported here). For the accuracy analyses, the responses were coded as binary values (1 = correct, 0 = incorrect), and we used the *glmer* function in the *lme4* package.

Latency analyses. Response times were faster for word targets than for nonword targets (640 vs. 741 ms, respectively; $t = 9.128$, $p < .001$). None of the planned contrasts approached significance: prime lexicality ($t = .004$, $p = .997$), word-prime frequency ($t = .454$, $p = .65$), and nonword-prime wordlikeness ($t = 1.291$, $p = .209$)—none of the interactions between target lexicality and prime type approached significance (all $ts < 1.485$, $ps > .137$).

Table 2

Mean correct response times (in ms) and accuracy (in brackets) for words and pseudowords in Experiment 1

	Prime type			
	High frequency word	Low frequency word	Pseudoword	Consonant string
Word	635 (96.7)	641 (97.1)	633 (96.9)	637 (95.3)
Pseudoword	734 (94.0)	730 (94.2)	733 (94.2)	735 (92.7)

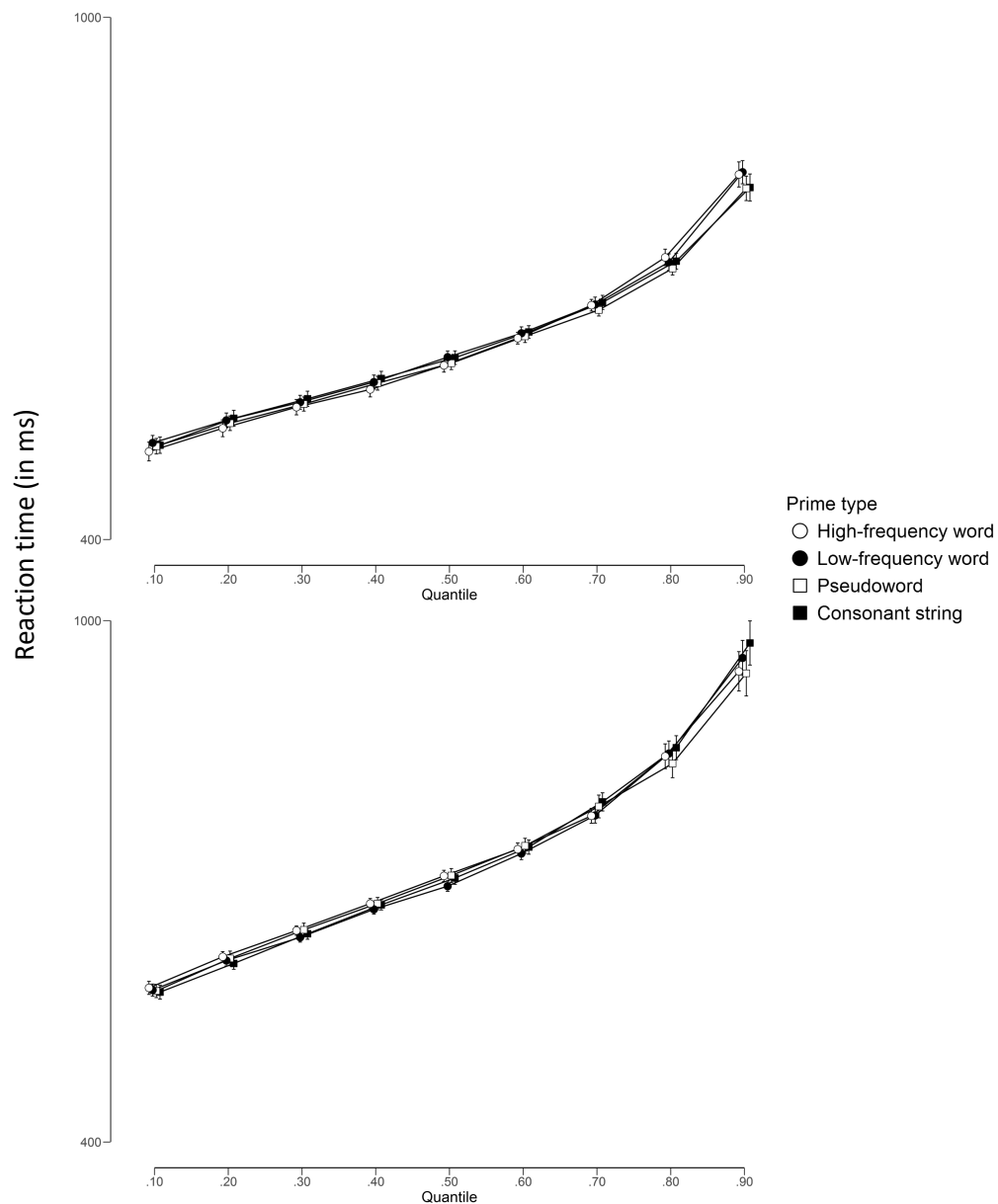
Accuracy analyses. Accuracy was higher for word targets than for nonword targets (93.8 vs. 96.5, respectively; $z = 6.949$, $p < .001$). Neither the effect of prime lexicality ($z = 1.778$, $p = .075$) nor the effect of prime word-frequency ($z = .634$, $p = .526$) was significant—the interactions with target lexicality were not significant either (all $ps > .524$). The effect of nonword-prime wordlikeness was

significant ($z = 2.896$, $p = .003$): participants were more accurate when the prime was a pseudoword than when the prime was a consonant string (95.6 vs. 94.0, respectively)—this effect occurred similarly for word and nonword targets, as can be deduced from the lack of interaction between the two factors ($z = .808$, $p = .419$).

To further examine whether the LME analyses missed some subtle effects (e.g., a facilitative effect in the leading edge of the RT distributions that could have been canceled by an inhibitory effect in the higher quantiles), we computed the .1, .2, .3, .4, .5, .6, .7, .8, and .9 quantiles for each participant and condition, and then calculated the values for each quantile across participants (i.e., vincentiled averages). As can be seen in the vincentile plots of the RT distributions shown in Figure 1, all priming conditions behaved similarly, thus corroborating the LME analyses.

Figure 1.

Group reaction time (RT) distributions for the four experimental conditions in Experiment 1 (upper panel: words; lower panel: nonwords) across quantiles (.1 to .9)



In sum, the present experiment, using orthographically legal nonwords as foils and a moderately large number of pairs in each condition (60 items/condition), showed that decision times to “yes” and “no” responses were not modulated by the wordlikeness of the unrelated primes, thus replicating previous experiments with wordlike nonword foils (Bailys et al., 2009; Loth & Davis, 2010a, Experiment 4; Norris & Kinoshita, 2008; Perea et al., 1998, 2010;

Perea et al., 2014, Experiment 1; Sereno, 1991). The only effect from the unrelated primes occurred in the accuracy data: accuracy was higher when the target stimulus was preceded by a pseudoword prime than when preceded by a prime composed of random consonants—note however that the size of the effect (while significant) was very small (95.6 vs. 94.0, respectively).

The BR model and IA model—under the selective inhibition hypothesis—can easily capture this pattern of findings. The question now is whether we could find an effect of the unrelated primes in masked priming lexical decision in an extreme scenario where the word/nonword discrimination is very easy. Experiment 2 was parallel to Experiment 1, except that we employed illegal nonwords as foils. As the rate of evidence accumulation for a “yes” or “no” response with illegal nonwords will be substantially higher than with wordlike nonwords, there is more room for a response congruency effect to appear. The idea is that participants are more likely to make an implicit decision based on the prime when the lexical status of prime and target is easily discriminable (e.g., the prime `gktbc` would produce evidence toward a “no” decision, whereas a prime like `house` would produce evidence toward a “yes” decision). Indeed, in a series of unpublished experiments, Loth and Davis (2010a) reported a response congruency effect in masked priming lexical decision when using unwordlike nonword targets (e.g., faster “yes” responses for `order-CATCH` than for `dqrki-CATCH`; faster “no” responses for `dqrki-SYKDD` than for `order-SYKDD`). A limitation of Loth and Davis (2010a) experiments, however, is that the two unrelated priming conditions (i.e., word primes vs. illegal primes) differed not only on lexical status (i.e., word vs. nonword), but also on their orthographic legality (legal vs. illegal).

Experiment 2

Method

Participants

Twenty-four new students from the same subject pool as in Experiment 1 participated in the experiment.

Materials

The set of stimuli was exactly the same as in Experiment 1 except for the nonword targets. Unlike Experiment 1, in which the foils were orthographically legal nonwords, in Experiment 2, all foils were illegal nonwords. All the 5-letter nonword foils contained an illegal trigram and a single vowel (e.g., BEFRM, TCFRO) so that they were hardly pronounceable—in a few cases (less than 9%), these nonwords had an orthographic neighbor (e.g., the nonword *chpja* has the very low-frequency word *chajá* [a Uruguayan dessert] as neighbor).

Procedure

The procedure was the same as in Experiment 1.

Results and discussion

As in Experiment 1, error responses and very short RTs (less than 250 ms: 1 response) were excluded from the latency analyses. The mean correct RTs and the accuracy in each condition are presented in Table 3. The statistical analyses were parallel to those reported in Experiment 1. There were 11,215 observations in the RT analyses.

Table 3

Mean correct response times (in ms) and accuracy (in brackets) for words and nonwords in Experiment 2

	Prime type			
	High frequency word	Low frequency word	Pseudoword	Consonant string
Word	542 (98.2)	539 (97.9)	537 (97.5)	565 (95.1)
Nonword	577 (97.5)	583 (97.5)	574 (97.2)	572 (97.5)

Latency analyses

Response times were faster for word targets than for nonword foils (546 vs. 576 ms, respectively; $t = 11.513$, $p < .001$).

Prime lexicality. While the effect of prime lexicality was not significant ($t = .758$, $p = .449$), we found a significant interaction of prime lexicality and target lexicality ($t = 6.585$, $p < .001$): lexical decision times on the target words were, on average, 10 ms faster when these were preceded by a word prime than when preceded by a nonword (pseudoword and consonant string) prime (541 vs. 551 ms, respectively; $t = 5.036$, $p < .001$), whereas lexical decision times on target nonwords were, on average, 7 ms faster when preceded by a nonword prime than when preceded by a word prime (573 vs. 580 ms, respectively; $t = 4.175$, $p < .001$).

Prime word-frequency. Neither the effect of the prime word-frequency ($t = .342$, $p = .732$), nor its interaction with target lexicality was significant ($t = 1.482$, $p = .138$).

Prime nonword-wordlikeness. The effect of prime nonword-wordlikeness was significant ($t = 3.573$, $p < .001$): response times were faster when the prime was a pseudoword than when the prime was a consonant string (556 vs. 568 ms, respectively). This effect interacted with target lexicality ($t = 8.229$, $p < .001$): lexical decision times on the target words were, on average, 28 ms faster when they were preceded by a pseudoword prime than when preceded by a consonant string (537 vs. 565 ms, respectively; $t = 7.704$, $p < .001$). In contrast, lexical decision times to target nonwords were, on average, 2 ms faster when the prime was a consonant string in comparison with pseudowords (572 vs. 574 ms, respectively; $t = 3.353$, $p < .001$). This latter difference, which was significant using inverse-transformed RT data, vanished when the analyses were conducted on the raw RT data, $t < 1$. As can be seen in the vincentile plot for nonword targets in Figure 2, there was a substantial advantage of the consonant string condition over the pseudoword condition in the leading edge of the RT distribution that vanished in the higher quantiles—note that inverse transformations put more weight in the short than in the long RTs. To examine this issue in greater depth, we conducted a 9 (Quantile) $\times$ 2 (Prime nonword-wordlikeness: pseudoword vs.

consonant string) analysis of variance. The main effect of Prime nonword-wordlikeness did not approach significance, $F < 1$, but critically, Prime nonword-wordlikeness interacted with Quantile, $F(8, 184) = 17.706$, $p < .001$: there was a substantial advantage of the consonant-string condition over the pseudoword condition in the leading edge of the RT distribution (29 ms in the .1 quantile [$p < .001$], 23 ms in the .2 quantile [$p < .001$] and 16 ms in the .3 quantile [$p = .009$], whereas the difference was (if anything) in the opposite direction at the highest quantile (-35 ms in the .9 quantile [$p = .108$])—we applied a Bonferroni correction in these simple effects tests. The nature of this interaction is discussed below. (Note that the parallel analyses with word targets revealed a prime nonword-wordlikeness effect [$F(1, 23) = 59.83$, $p < .001$] that was approximately the same size across the quantiles [interaction: $F(8, 184) = 1.03$, $p = .41$].)

Accuracy analyses

We did not find any differences in accuracy between the words and nonwords ($z = .09$, $p = .926$).

Prime lexicality. The prime lexicality effect was significant ($z = 2.800$, $p = .005$): participants were more accurate when the prime was a word than when the prime was a nonword (97.8 vs. 96.8, respectively). This effect interacted with target lexicality ($z = 2.450$, $p = .014$): For word targets, responses were more accurate when the unrelated prime was a word than when it was not a word (98.1 vs. 96.4; $z = 3.678$, $p < .001$), whereas this effect did not occur for nonword targets ($z = .236$, $p = .813$).

Prime word-frequency. Neither the effect of prime word-frequency ($z = .5$, $p = .619$) nor its interaction with target lexicality approached significance ($z = .34$, $p = .735$).

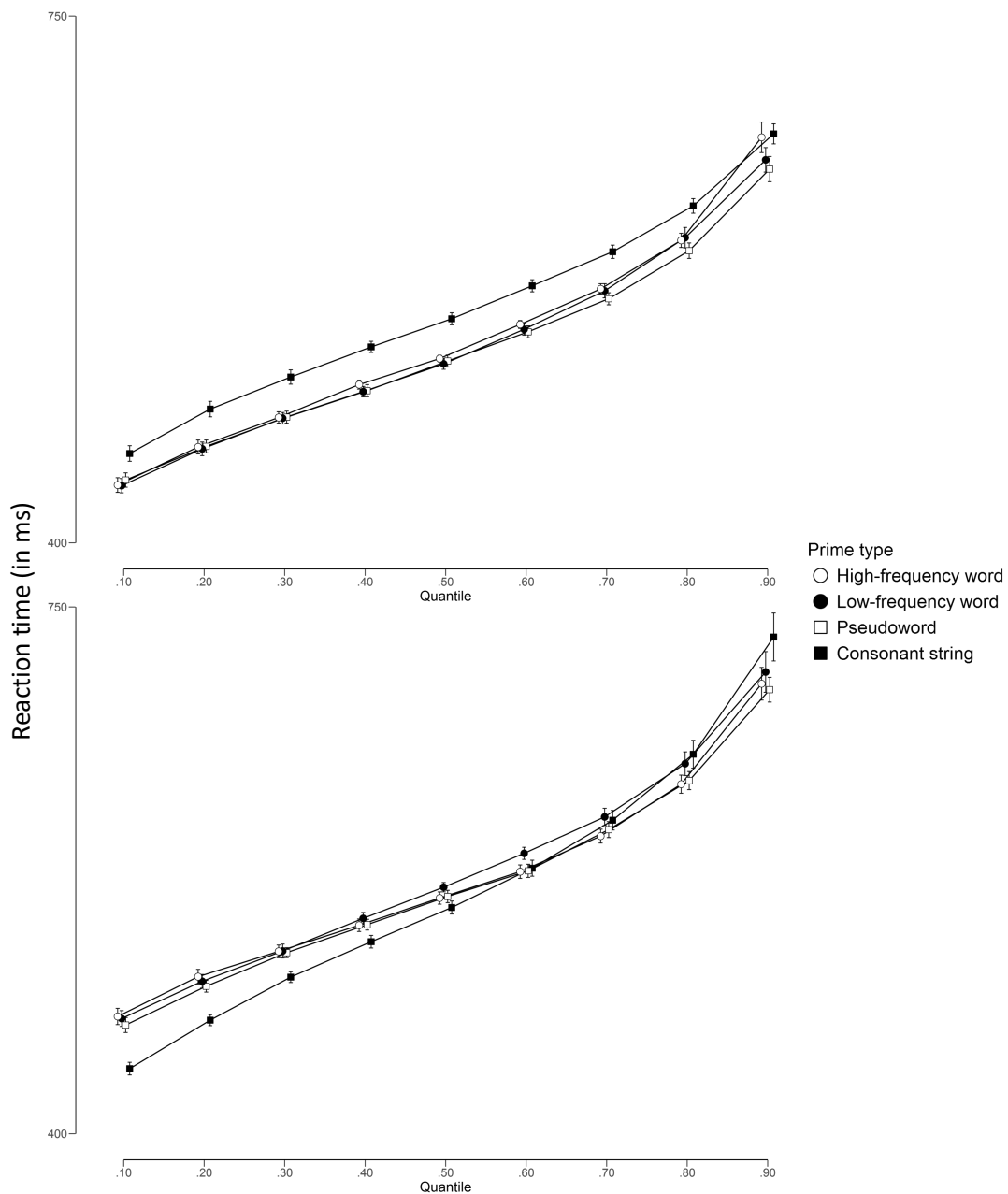
Prime nonword-wordlikeness. The effect of prime nonword-wordlikeness approximated significance ($z = 1.95$, $p = .052$). This effect interacted with target lexicality ($z = 2.660$, $p = .007$): for word targets, participants were more accurate when the prime was a pseudoword than when it was consonant string (97.6 vs. 95.1; $z = 3.482$, $p < .001$), but this difference did not approach significance for nonword targets ($z = .508$, $p = .611$). As shown in the vincentile plot displayed in

Figure 2 (top panel), the RT distributions for the word targets were remarkably similar when preceded by an unrelated high-frequency word prime, an unrelated low-frequency word prime, or a pseudoword prime (i.e., the orthographically legal primes), whereas the RT distribution for the word targets preceded by a consonant-string prime showed a location shift to the right. Thus, the effect from consonant-string primes is not a simple response congruency effect (i.e., both consonant-strings primes and pseudoword primes presumably provide evidence for “no” responses in lexical decision, but the pseudoword primes behaved similarly to the word primes). Instead, what this pattern of data suggests is that participants considered the orthographic characteristics of the primes regardless of their lexical status. This occurred fast enough so that an orthographically illegal prime produced evidence toward a “no” decision, whereas an orthographically legal prime produced evidence toward a “yes” decision²¹. This reinforces the view that the participants in the experiment 2 could have evaluated the legality of the stimuli rather than accessing the mental lexicon (i.e., a judgment legality task). Indeed, Forster, Mohan, and Hector (2003) showed that word-frequency effects (i.e., a marker of lexical access) are dramatically diminished in lexical decision experiments with illegal nonword foils—consistent with this view, the lexicality effect in the current experiment was dramatically smaller than in Experiment 1 (31 vs. 97 ms, respectively).

²¹ We thank the reviewers and the editor for suggesting this explanation.

Figure 2

Group reaction time (RT) distributions for the four experimental conditions in Experiment 2 (upper panel: words; lower panel: nonwords) across quantiles (.1 to .9).



As shown in the vincentile plot displayed in Figure 2 (top panel), the RT distributions for the word targets were remarkably similar when preceded by an unrelated high-frequency word prime, an unrelated low-frequency word prime, or a pseudoword prime (i.e., the orthographically legal primes), whereas the RT distribution for the word targets preceded by a consonant-string prime showed a

location shift to the right. Thus, the effect from consonant-string primes is not a simple response congruency effect (i.e., both consonant-strings primes and pseudoword primes presumably provide evidence for “no” responses in lexical decision, but the pseudoword primes behaved similarly to the word primes). Instead, what this pattern of data suggests is that participants considered the orthographic characteristics of the primes regardless of their lexical status. This occurred fast enough so that an orthographically illegal prime produced evidence toward a “no” decision, whereas an orthographically legal prime produced evidence toward a “yes” decision²². This reinforces the view that the participants in the experiment 2 could have evaluated the legality of the stimuli rather than accessing the mental lexicon (i.e., a judgment legality task). Indeed, Forster, Mohan, and Hector (2003) showed that word-frequency effects (i.e., a marker of lexical access) are dramatically diminished in lexical decision experiments with illegal nonword foils—consistent with this view, the lexicality effect in the current experiment was dramatically smaller than in Experiment 1 (31 vs. 97 ms, respectively).

For the nonword targets, the RT distributions were again very similar when preceded by a high-frequency unrelated prime, a low-frequency unrelated prime, or a pseudoword prime (i.e., the orthographic legal primes). There is, however, a somewhat unusual pattern in the RT distribution of the consonant string condition: we found a large legality congruency effect in the initial quantiles of the RT distribution (i.e., shorter “no” response times to nonword targets when the prime was orthographically illegal than when the prime was orthographically legal) that vanished in the higher quantiles. This pattern is puzzling, as one would expect an effect to increase its size across quantiles (e.g., word-frequency effect; Ratcliff et al., 2004) or to be approximately the same size across quantiles (e.g., masked identity priming; Gomez et al., 2013). The existence of a large effect in the leading edge of the RT distribution that disappears in the higher quantiles has been documented in a numerical categorization task (“is the digit higher than 5?”) with masked primes (e.g., prime: 7; target: 9 vs. prime: 3; target: 9; see Kinoshita &

²² We thank the reviewers and the editor for suggesting this explanation.

Hunt, 2008). Kinoshita and Hunt (2008) posited the existence of an “inhibitory control mechanism” that rejects an activated response when the cognitive system detects that “the (supraliminal) target is not the source of the response” (p. 1331-1332). However, it is unclear how a prime can be discounted when it is not visible—the theories on prime discounting originated from studies using visible primes (e.g., see Huber & O’Reilly, 2003). Indeed, de Wit and Kinoshita (2014) acknowledged that “in masked priming, no prime discounting occurs because the prime is masked and hence the participants are unaware that the evidence accumulated from the prime comes from a difference source” (p. 1739). Clearly, the disappearance of the legality congruency effect in the higher quantiles for the nonword targets grants some explanation. The following is an admittedly ad hoc account that should be explored in more targeted studies.

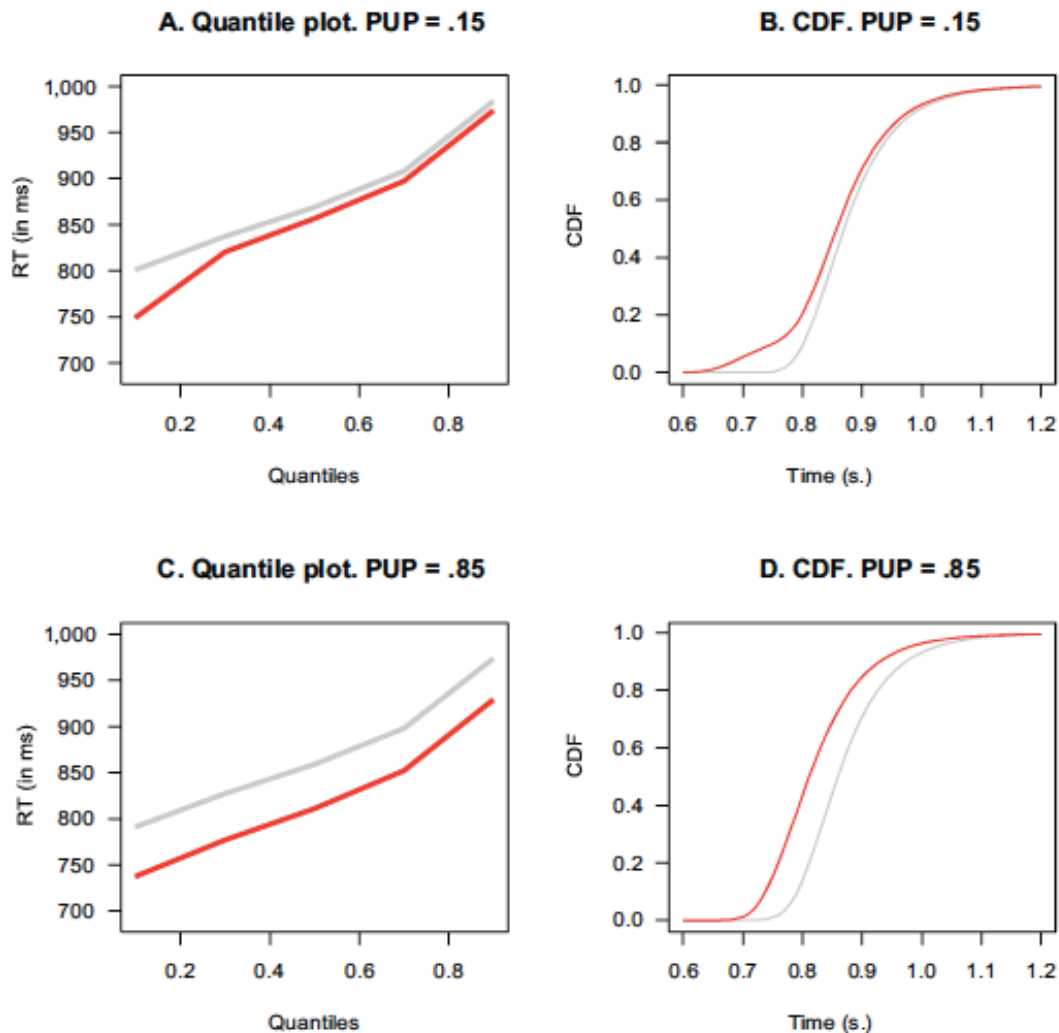
Here, we propose a plausible mechanism within the structure of evidence accumulation modeling²³. We term it the “probabilistic use of the prime” (PUP) hypothesis. The intuition is straightforward: suppose that the observer does not use the information provided by the prime in every trial. Instead, the prime-target relationship is utilized in a proportion of trials (i.e., there is a probability greater than 0 and less than 1 that in a given trial the information provided by the prime is used). This hypothesis can be implemented in many different ways. However, given prior research on masked priming effects, we believe that assuming that prime-target congruency affects the time of encoding+response (T_{er} parameter) is reasonable (see Gomez et al., 2013, for discussion). (Note that the diffusion model cannot disentangle the time of encoding from the time of response). If the probability of using the prime-target relationship were 1 ($PUP = 1$), then the priming effect would produce a location shift in the RT distribution. Critically, as this probability goes down, the effect on the tail is progressively reduced (see Figure 3). Note that at relatively high PUPs (.85 in the panels C and D of Figure 3), there seems to be a location shift in the RT distributions. This matches the results with word targets in the current experiment and also the masked identity

²³ We are indebted to Pablo Gomez for suggesting this explanation. While we use terms of the diffusion model account of lexical decision (Gomez et al., 2013; Ratcliff et al., 2004), Norris (2009) showed that this model can be subsumed within the Bayesian Reader model.

priming effects reported by Gomez et al. (2013; see also Perea, Marcet, Lozano, & Gomez, 2018). Importantly, when the PUP is low (.15 in the panels A and B of Figure 3), we can observe a pattern very similar to that found with nonword targets. That is, if prime-target legality congruency affects the time of encoding+response (T_{er} parameter) in a relatively small proportion of trials (e.g., via an implicit decision to the prime stimulus), the model predicts an advantage of the congruent condition in the leading edge of the RT distribution that vanishes in higher quantiles. Why does the pattern of effects differ for word and nonword targets? We could speculate that the PUP might be related to the usefulness of the prime. If we assume that participants were basing their implicit decision on prime legality rather than prime lexicality, three out of the four priming conditions were congruent (i.e., high-frequency words, low-frequency words, orthographically legal pseudowords). Thus, in a large majority of trials, the orthographic legality of the primes would be helpful to reach a “yes” decision (i.e., the implicit decision to the prime matches the decision to the target). If we assume that this affects mainly the T_{er} parameter of the model (i.e., time of encoding+response), one would obtain a legality congruency effect of similar size across the quantiles, as actually occurred. In contrast, for nonword targets, only one out of four priming conditions was congruent. In this scenario, prime-target congruency would be substantially less helpful to reach a “no” decision and it may have been used in only a limited proportion of trials. While admittedly ad hoc, the PUP hypothesis offers a reasonable account of the RT distributions for word and nonword targets in the present experiment (i.e., compare Figures 2 and 3). Importantly, the diminishing masked priming effect in the higher quantiles obtained by Kinoshita and Hunt (2008; Experiment 1) in a numerical categorization task can also be accommodated by the PUP hypothesis: in their experiment, only one out of four priming conditions were congruent with the response.

Figure 3

Quantile plots (left) and cumulative density functions (CDFs; right) for two conditions that differ in the probabilistic use of the prime (PUP): low proportion of trials (PUP .15; top panels) versus low proportion of trials (PUP .85, bottom panels).



Note. Prime-target relationship is assumed to affect the Time of Encoding + Response (Ter parameter) of the target. RT = Reaction Time.

In sum, the present experiment represents a demonstration of a response congruency effect in masked priming when the participants' task requires the discrimination of words and illegal nonwords (e.g., *bunte*-JAULA [the English for *funt*-CAGE] produced faster "yes" responses than *pdtrv*-JAULA [*pdtr*-CAGE]; *ghdk*-BEFRM produced faster "no" responses than *fiok*-BEFRM). Importantly, these differences were completely absent with orthographically legal primes (i.e., high-frequency words, low-frequency words, and pseudowords

behaved similarly; see Figure 2) and only occurred for consonant-string primes. This finding replicates and extends Loth and Davis' (2010a) Experiment 1. In Loth and Davis' Experiment 1, participants responses were 23 ms faster when the target word were preceded by a congruent prime (i.e., a high-frequency word; 475 ms) than when preceded by an incongruent prime (i.e., an illegal nonword; 498 ms)—we also found a 23-ms difference when comparing these two conditions.

What are the implications of this congruency effect for models of visual word recognition? While computational models of visual word recognition have been used to simulate lexical decision experiments with illegal nonwords as foils (e.g., Loth & Davis, 2010b; Norris, 2009; Ratcliff et al., 2004), one might argue that under these circumstances, the participants' task is not lexical decision per se, but rather an orthographic legality task (see Forster et al., 2003, for discussion)—note that we did not find a simple response congruency effect (i.e., unrelated pseudoword primes behaved just as unrelated word primes), but a legality congruency effect. To examine the involvement of lexical processing in Experiment 2, we calculated the Pearson correlation between the mean RT of each word and its Zipf frequency. The idea is that if participants are treating the task as an orthographic legality task, frequency effects should be very small. This analysis showed a negligible correlation between these two variables ($r = -0.04$, $p = .55$), thus supporting the view that the word/nonword task in the current experiment did not involve lexical access. (The parallel correlation coefficient in Experiment 1 was $r = -0.27$, $p < .001$.) Thus, we believe that it would be premature to make claims on whether the legality congruency effect in the current experiment poses serious problems for the masked priming account of the BR model or the selective inhibition hypothesis of the IA model in a standard lexical decision task—note that these two accounts would predict similar response times for all priming conditions.

To create a more constraining scenario for the BR and IA models in an easy-to-perform lexical decision task without the risk of altering its lexical nature, we designed Experiment 3. Experiment 3 was similar to Experiment 2 in which the word/nonword discrimination was easy: the nonword foils did not have any

close neighbors. Critically, all nonword foils were orthographically legal (GAZUR). In this scenario, lexical decisions need to be made in terms of lexical activation rather than on orthographic legality, so that the BR and the IA model—under the selective inhibition hypothesis— would predict a null effect of response congruency.

Experiment 3

Method

Participants

Twenty-four new students from the same subject pool as in Experiment 1 and 2 participated in the experiment.

Materials

The set of stimuli was exactly the same as in Experiment 1 and 2 except for the nonword targets. Unlike Experiment 2—where foils were illegal nonwords—in Experiment 3 all foils were orthographically legal nonwords. We created the foils by changing two letters from Spanish words using Pseudo software (van Heuven, 2002). Then, we filtered the output by extracting only those nonwords with no close lexical neighbors (i.e., no one-letter substitution-letter neighbors; no transposed-letter neighbors; no addition-letter neighbors, and no deletion-letter neighbors). This produced a list of more than 300 orthographically legal hermit nonwords (e.g., LEDUL, RUPUO) and we randomly selected 240 nonwords from this set.

Procedure

It was the same as in Experiments 1 and 2.

Results and discussion

As in Experiments 1 and 2, error responses and very short RTs (less than 250 ms: 1 response) were omitted from the latency analyses. The mean correct RTs and the accuracy in each experimental condition are presented in Table 4. The

statistical analyses were parallel to those reported in the previous experiments. There were 11,167 observations in the RT analyses.

Table 4

Mean correct response times (in ms) and accuracy (in brackets) for words and nonwords in Experiment 3

	Prime type			
	High frequency word	Low frequency word	Pseudoword	Consonant string
Word	644 (96.7)	640 (96.9)	643 (95.6)	647 (95.6)
Pseudoword	701 (97.5)	708 (98.2)	698 (97.9)	709 (97.0)

Latency analyses

Response times were, on average, 60 ms faster for word targets than for nonword targets (644 vs. 704 ms, respectively; $t = 5.775$, $p < .001$).

Prime lexicality. Neither the effect of prime lexicality nor its interaction with target type was significant (both $ts < 1.26$, $ps > .21$). Indeed, lexical decision times were only 1 ms faster when the prime was a word than when the prime was a nonword (673 vs. 674 ms, respectively).

Prime word-frequency. Neither the effect of word-prime frequency nor its interaction with target type approached significant (both $ts < 1$, $ps > .65$). Participants' responses were 1 ms faster when the target was preceded by a high-frequency word than when preceded by a low-frequency word (673 vs. 674 ms, respectively).

Prime nonword-wordlikeness. Neither the effect of prime nonword-wordlikeness ($t = 1.941$, $p = .065$) nor its interaction with target type were significant ($t = .403$, $p = .687$)—note that, for word targets, the difference between the pseudoword and consonant string conditions was only 4 ms (643 vs. 647 ms, respectively; $t = 1.275$, $p = .208$).

Accuracy analyses

Accuracy was higher for nonword targets than for word targets (97.7 vs. 96.2, respectively; $z = 4.44$, $p < .001$).

Prime lexicality. The effect of prime lexicality was significant ($z = 2.264$, $p = .02$): participants were more accurate when the prime was a word than when the prime was a nonword (97.3 vs. 96.5, respectively)—this effect occurred similarly for word and nonword targets, as can be deduced from the lack of interaction between prime and target lexicality ($z = .789$, $p = .429$).

Prime word-frequency. Neither the effect of word prime frequency nor its interaction with target lexicality was significant (both z s < 1 , both p s $> .42$). Subjects were only 0.5% more accurate when the prime was a low-frequency word than when the prime was a high-frequency word (97.6 vs. 97.1, respectively).

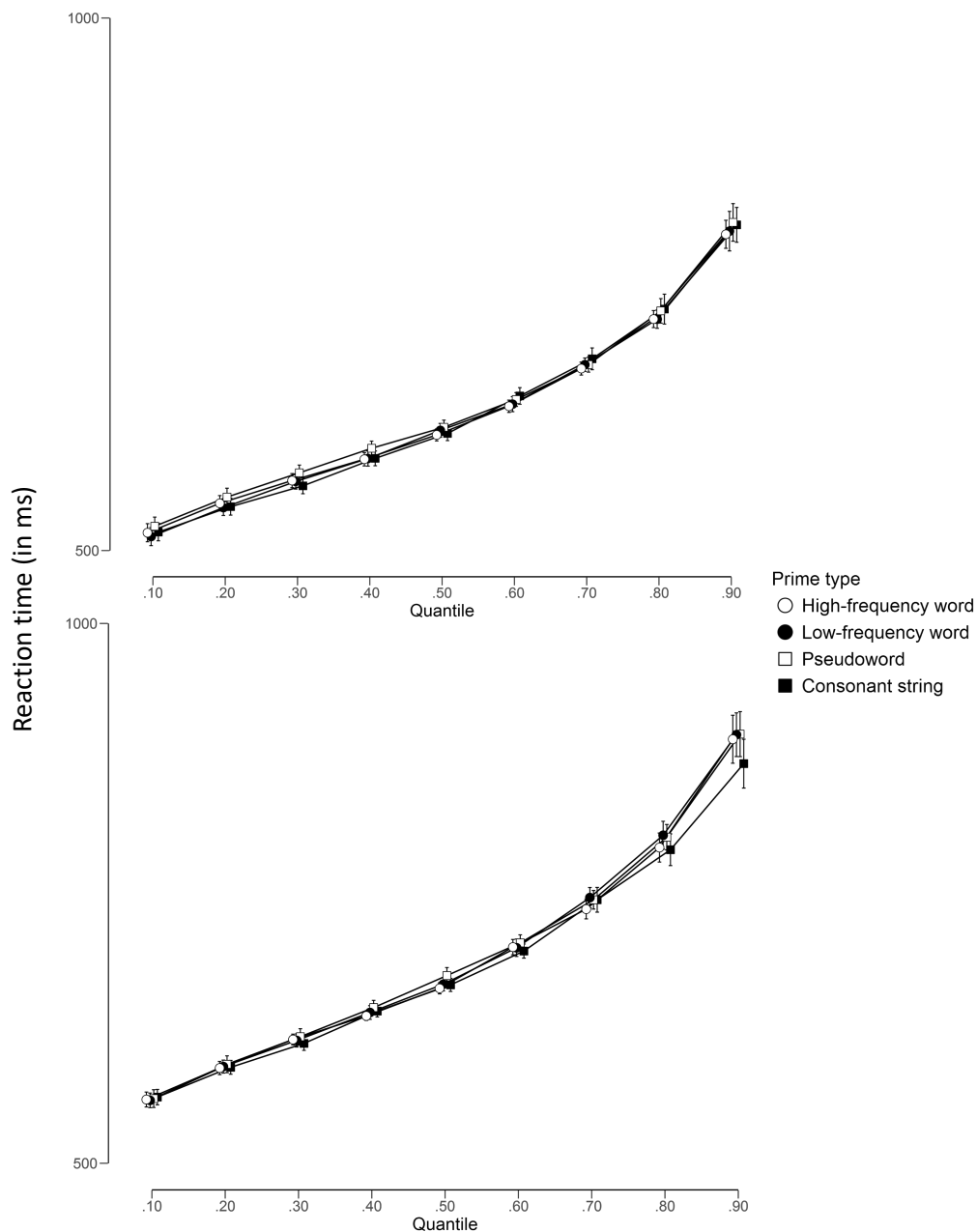
Prime nonword-wordlikeness. Neither the effect of prime nonword-wordlikeness nor its interaction with target type approached significance (both z s < 1.24 , p s $> .21$). The difference in accuracy between pseudoword and consonant string primes is only 0.5% (96.8 vs. 96.3, respectively).

As can be seen in the vincentile plots for both the word and nonword targets displayed in Figure 4, all priming conditions behaved similarly, thus corroborating the LME analyses.

Therefore, the present experiment showed that, when using orthographically legal hermit nonwords as foils, lexical decision times on the target stimuli were similar regardless of the characteristics of the unrelated primes. This pattern of data is similar to that found in Experiment 1—which employed orthographically legal nonwords matched with words in sublexical characteristics. Furthermore, the Pearson correlation between the mean RT of each word and its Zipf frequency in the current experiment was similar to that found in Experiment 1 ($r = -0.28$, $p < .001$ and $r = -0.27$, $p < .001$, respectively), thus showing that the two experiments involve lexical access. The only effect from unrelated primes occurred in the accuracy data: accuracy was 0.8% higher when the target stimuli were preceded by a word prime than when preceded by a nonword prime (97.3 vs. 96.5, respectively).

Figure 4

Group reaction time (RT) distributions for the four experimental conditions in Experiment 3 (upper panel: words; lower panel: nonwords) across quantiles (.1 to .9)



The null finding in the latency data may be in contrast to Loth and Davis' (2010a) Experiment 2, who found that lexical decision times for word targets were faster when preceded by a high-frequency word prime than when preceded by an illegal nonword prime when the foils were pronounceable nonwords with no neighbors (e.g., TEIR, DULEW). However, the effect size in Loth and Davis'

(2010a) Experiment 2 was only 8 ms—the parallel difference in the current experiment was 3 ms (see Table 4). Thus, we prefer to remain cautious about this small effect.

General Discussion

The present experiments were designed to examine in detail whether the lexical status and wordlikeness of unrelated primes (high-frequency words, low-frequency words, orthographically legal pseudowords, and random consonant strings) could affect target processing in masked priming lexical decision. We did so in three different scenarios that varied in the difficulty of the word/nonword discrimination: orthographically legal nonwords matched in subsyllabic elements with the target words (Experiment 1), orthographically illegal nonword foils (Experiment 2), and orthographically legal nonwords with no neighbors (Experiment 3). When using orthographically legal nonwords, response times to words and nonwords did not differ as a function of the wordlikeness of the prime stimuli (Experiments 1 and 3; see Figures 1 and 4), thus extending previous research with orthographically legal foils (see Table 1 for a summary) with a more thorough set of conditions. We only found an effect of the unrelated primes when the nonwords were orthographically illegal (Experiment 2) and restricted to the primes composed of consonant strings (see Figure 2). We now discuss the implications of these findings for leading models of visual word recognition.

The Bayesian Reader (BR) model predicts similar response times for words preceded by an unrelated prime regardless of its wordlikeness in lexical decision (see Loth & Davis, 2010b; Norris & Kinoshita, 2008, for simulations). Consistent with the BR model, all unrelated priming conditions behaved similarly in standard word recognition scenarios (i.e., when the foils were orthographically legal, as in Experiments 1 and 3).²⁴ However, when the word/nonword

²⁴ The only minor differences across conditions were in the accuracy analysis. In Experiment 1 there was a 1.6% increase in the error rates when the primes were composed of random consonants relative to pseudoword primes. In Experiment 3 accuracy was 0.8% higher when the target stimuli were preceded by a word prime than when preceded by a nonword prime.

discrimination was much easier (i.e., using illegal nonwords as foils [e.g., *BEFRM*, *TCFRO*] instead of orthographically legal nonwords), we found an effect from the unrelated primes: word response times were longer when the masked prime was a consonant string than when the masked prime was an orthographically legal nonword (*ghdk-TRUE* produced longer response times than *fiok-TRUE*; see Figure 2); and nonword responses were faster when preceded by a consonant-string prime than when preceded by a orthographically legal prime, but only in the initial quantiles of RT distributions (*ghdk-BEFRM* produced faster responses than *fiok-BEFRM*; see Figure 2). The question now is whether this latter finding poses some problems for the BR account of masked priming lexical decision. One might argue that the task in Experiment 2 may have switched from lexical decision to an orthographic legality decision. In an orthographic legality decision task, the BR model could quickly accumulate evidence for “yes” or “no” responses from the masked primes: as all the nonword foils violated the orthographic constraints, then anything legal must be a word. That is, participants could carry over from prime to target the estimate of the probability that the visual input is orthographically well formed, not the lexical status. As a result, the BR model could predict slower “yes” responses when the prime is composed of random consonants than when the prime is an orthographically legal stimulus (e.g., a pseudoword or a word), as actually occurred in Experiment 2. Conversely, the model could predict faster “no” responses when the prime is a consonant string than when the prime is an orthographically legal stimulus (e.g., a pseudoword or a word). Indeed, as indicated in the Introduction, the BR model can readily accommodate response congruency effects in two-choice tasks other than lexical decision (e.g., semantic categorization; see de Wit & Kinoshita, 2014). Thus, the present findings offer some empirical support to the BR account of masked priming in standard lexical decision experiments—they are also a demonstration of how the nature of the decision required by the task modulates masked priming effects.

The interactive activation model (McClelland & Rumelhart, 1981) offers different predictions depending on whether there is homogeneous or selective lateral inhibition (see Bailyss et al., 2009, for simulations). Clearly, the

homogeneous lateral inhibition hypothesis cannot capture the pattern of effects in any of the three experiments (i.e., this hypothesis wrongly predicts longer “yes” decisions in the high-frequency word priming condition than in the nonword priming conditions). Instead, the pattern of the data with orthographically legal nonwords (Experiments 1 and 3) favors Davis and Lupker’s (2006) selective lateral inhibition hypothesis: the processing of the word `TRUE` is not modified by the wordlikeness of the unrelated prime, whether it is a high-frequency word like *food* or a consonant string like *ghdk*. The only potential drawback is that response times are faster for *food-TRUE* than for *ghdk-TRUE* when the foils were orthographically illegal (Experiment 2). While the original IA model under the selective lateral inhibition hypothesis cannot predict this difference, one might argue that participants in Experiment 2 were making an orthographic legality judgment rather than a lexical decision. This orthographic legality judgment task would be beyond the scope of the IA model. A somewhat crude approximation would be to use “overall lexical activity” as a source of evidence to speed up responses in this scenario, as in the “fast-guess” criterion proposed by the MROM (Grainger & Jacobs, 1996) or the SCM (see Davis, 2010). If the fast-guess criterion were set low when the nonword foils do not generate practically any lexical activity, those words whose primes produce some lexical activation would enjoy some advantage relative to those words whose primes barely produced any lexical activation. Indeed, Loth and Davis (2010b) conducted simulations that showed that, depending on the value of the parameters, the SCM could predict an effect of response congruency (e.g., *fiok-TRUE* faster than *ghdk-TRUE*) when the nonword foils are unwordlike and orthographically illegal. However, it is unclear to us how this fast-guess mechanism can accommodate the findings of Experiment 3 with hermit nonwords as foils—note that these stimuli also generate a very low level of activation in the lexicon. Further empirical and computational work is necessary to examine the similarities and differences on word vs. nonword responses in lexical decision vs. orthographic judgment tasks.

In summary, the present experiments were designed to examine whether response congruency effects with wordlike/unwordlike unrelated primes could be found in lexical decision. When the foils were orthographically legal nonwords,

both wordlike and unwordlike priming conditions behaved similarly (Experiments 1 and 3). These findings suggest that participants are unable to establish the lexical status of the masked prime or even to determine the likelihood of its being a word in a standard lexical decision task, thus providing empirical support to the BR model and the selective inhibition hypothesis of the IA model. Furthermore, our findings are a demonstration that detecting invariances (i.e., the null effect of lexical status of unrelated primes on target processing in masked priming lexical decision) is important to progress in science (see Rouder et al., 2009). Thus, at a methodological level, the current experiments are useful to choose the appropriate baseline in masked priming lexical decision experiments. We only found an effect of the unrelated primes on target processing when the foils were orthographically illegal and the primes were composed of consonant strings. However, in this scenario, one might argue that the participants' responses could have been driven by orthographic properties rather than lexical access. Further experimentation using measures with a better temporal resolution (e.g., ERPs) may offer important insights on the time course of the effects of response congruency across tasks.

Unveiling the boost in the sandwich priming technique

Abstract

The masked priming technique (e.g., #####-house-HOUSE) is the gold-standard tool to examine the initial moments of word processing. Lupker and Davis (2009) showed that adding a pre-prime identical to the target produced greater priming effects in the sandwich technique (which compares #####-HOUSE-house-HOUSE vs #####-HOUSE-fight-HOUSE). While there is consensus that the sandwich technique magnifies the size of priming effects relative to the standard procedure, the mechanisms underlying this boost are not well understood (i.e., does it reflect quantitative or qualitative changes?). To fully characterize the sandwich technique, we compared the sandwich and standard techniques by examining the RTs and their distributional features (delta plots; conditional-accuracy functions), comparing identity vs. unrelated primes. Results showed that the locus of the boost in the sandwich technique was two-fold: faster responses in the identity condition (via a shift in the RT distributions) and slower responses in the unrelated condition. We discuss the theoretical and methodological implications of these findings.

Key words: masked priming; lexical decision; visual word recognition; interactive activation

Introduction

Forster and Davis' (1984) masked priming technique has become the gold-standard tool for researchers interested in the initial moments of letter and word processing. This technique allows researchers to manipulate the relationship between a forwardly masked and briefly presented prime (approximately 30-50 ms) and a target item, while minimizing the non-automatic processes induced by visible primes (see Grainger, 2008, for a review). The type of prime-target relationship depends on the ingenuity of the researchers' working hypothesis (e.g., to cite three instances: phonological priming: #####-hint-MINT vs. #####-pint-MINT; translation priming: #####- maison-HOUSE vs. #####-appel-HOUSE; identity priming: #####-house-HOUSE vs. #####-fight-HOUSE). Proof of the ubiquity of this technique is that it has been employed in a wide range of approaches, from behavioral research—with response times (henceforth, RTs) and accuracy as dependent variables (e.g., Forster et al., 1987; Grainger & Ferrand, 1994)—to computational modeling (e.g., Gomez et al., 2013) and cognitive neuroscience (e.g., functional magnetic resonance imaging [fMRI]: Dehaene et al., 2001; event related potentials [ERPs]: Grainger & Holcomb, 2009; magnetoencephalography [MEG]: Monahan et al., 2008).

Because of the short stimulus-onset asynchrony between prime and target, the magnitude of masked priming effects is small, in the range of 10-50 ms in behavioral studies. Thus, a problem faced by researchers using this technique is that it may not be sensitive enough to detect subtle manipulations or interactions (e.g., lght-LIGHT vs. liht-LIGHT; see Grainger, 2008). This issue is more prominent when designing experiments with developing readers or special populations (e.g., readers with dyslexia or deaf readers), for whom there tends to be more within-participant variability in the RTs compared to university students (i.e., the typical participants recruited in research studies; see Lété & Fayol, 2013). One approach to overcome this intrinsic limitation of masked priming is statistical, by conducting over-powered large-scale experiments (e.g., Adelman et al.'s, 2014, mega-study with more than 1,000 participants, each responding to 840 stimuli). While highly valuable, this approach is costly in

resources and in time, and it does not fully address the within-participant variability.

Another approach is to modify the technique to magnify the effects; this was a basic idea underlying the inception of the Lupker and Davis' (2009) sandwich technique. The sandwich technique consists of a modification of the standard masked priming paradigm. The forward masked is followed by an initial masked prime identical to the target (the so-called pre-prime), followed by the prime of interest (see Figure 1 for an example). As in the standard technique, the target replaces the prime and is displayed until the participant responds. Note that both methods involve a prime that is pre- and post-masked, so strictly speaking, both procedures are forms of masked priming; hence, in the remainder of the article, we will refer to them as *standard* and *sandwich* techniques.

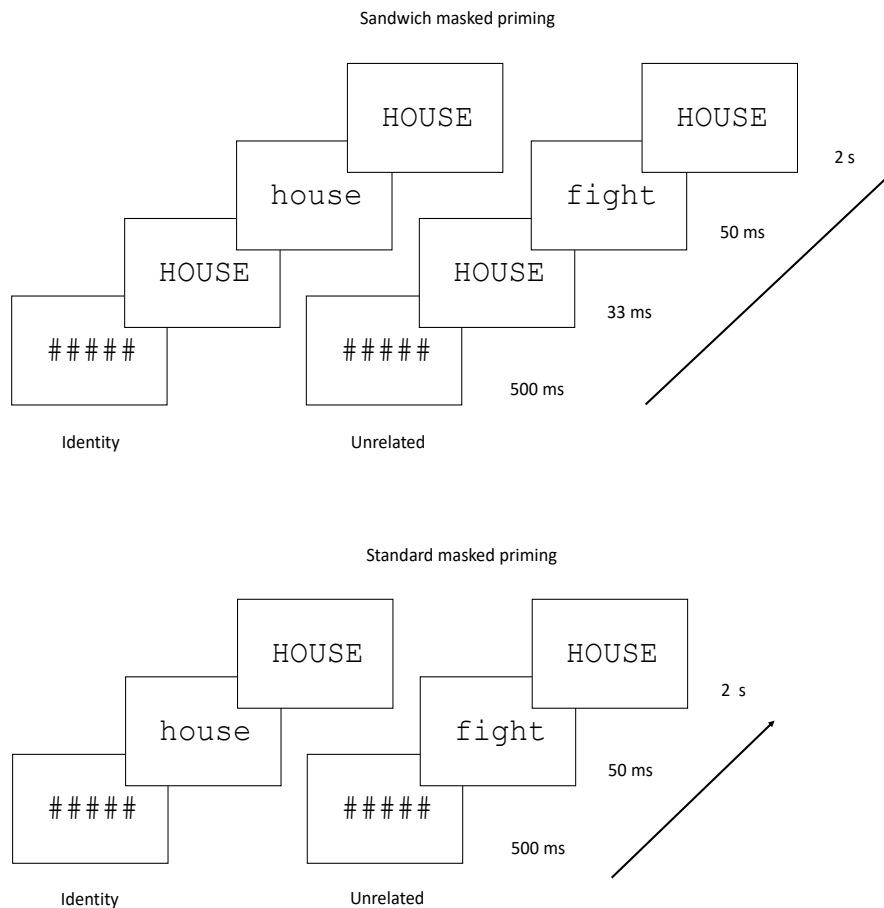
As first shown in a series of form-priming experiments conducted by Lupker and Davis (2009), the sandwich technique produces greater priming effects than the standard technique. This boost has since been replicated with various types of prime-target relationships (e.g., see Comesaña et al., 2015; Perea et al., 2014; Stinchcombe et al., 2012; Trifonova & Adelman, 2018, for behavioral evidence, and Ktori et al., 2012, for electrophysiological evidence). Indeed, because of the larger priming effects, the sandwich methodology has become the preferred option in many recent masked priming studies (e.g., Campos et al., 2020; Colombo, et a., 2020; Gutiérrez-Sigut et al., 2017; Ktori et al., 2014; Lupker et al., 2020a, 2020b; Lupker et al., 2015; Meade, Grainger et al., 2020; Ziegler et al., 2014; Taikh & Lupker, 2020).

As first shown in a series of form-priming experiments conducted by Lupker and Davis (2009), the sandwich technique produces greater priming effects than the standard technique. This boost has since been replicated with various types of prime-target relationships (e.g., see Comesaña et al., 2015; Perea et al., 2014; Stinchcombe et al., 2012; Trifonova & Adelman, 2018, for behavioral evidence, and Ktori et al., 2012, for electrophysiological evidence). Indeed, because of the larger priming effects, the sandwich methodology has become the preferred option in many recent masked priming studies (e.g.,

Campos et al., 2020; Colombo, et al., 2020; Gutiérrez-Sigut et al., 2017; Ktori et al., 2014; Lupker et al., 2020a, 2020b; Lupker et al., 2015; Meade, Grainger et al., 2020; Ziegler et al., 2014; Taikh & Lupker, 2020).

Figure 1

Sequence of events with the sandwich masked priming technique (top panel) and the standard masked priming technique (bottom panel)



Notwithstanding its utility in amplifying priming effects, the sandwich technique requires some further examination. The specific mechanisms at play and how they differ from the standard technique, have not been fully elucidated (see Lupker & Davis, 2009 vs. Trifonova & Adelman, 2018). Thus, it is not yet well understood if the priming effects in the sandwich technique are simply magnified or, instead, the underlying processes are fundamentally altered. Exploring this question is the main goal of the present work. In the original Lupker and Davis (2009) experiments, task procedure (sandwich vs. standard) was manipulated between subjects, making it difficult to compare the nuances of the

two techniques. Trifonova and Adelman (2018) have been the only researchers to date that employed a within-subject design for task procedure. However, their form-priming experiments focused on whether the characteristics of the pre-prime influenced word processing (e.g., how the pre-prime `PROTECT` influences the processing of the target `PROJECT`) rather than a direct comparison of the two techniques using their most common setup as in Figure 1.

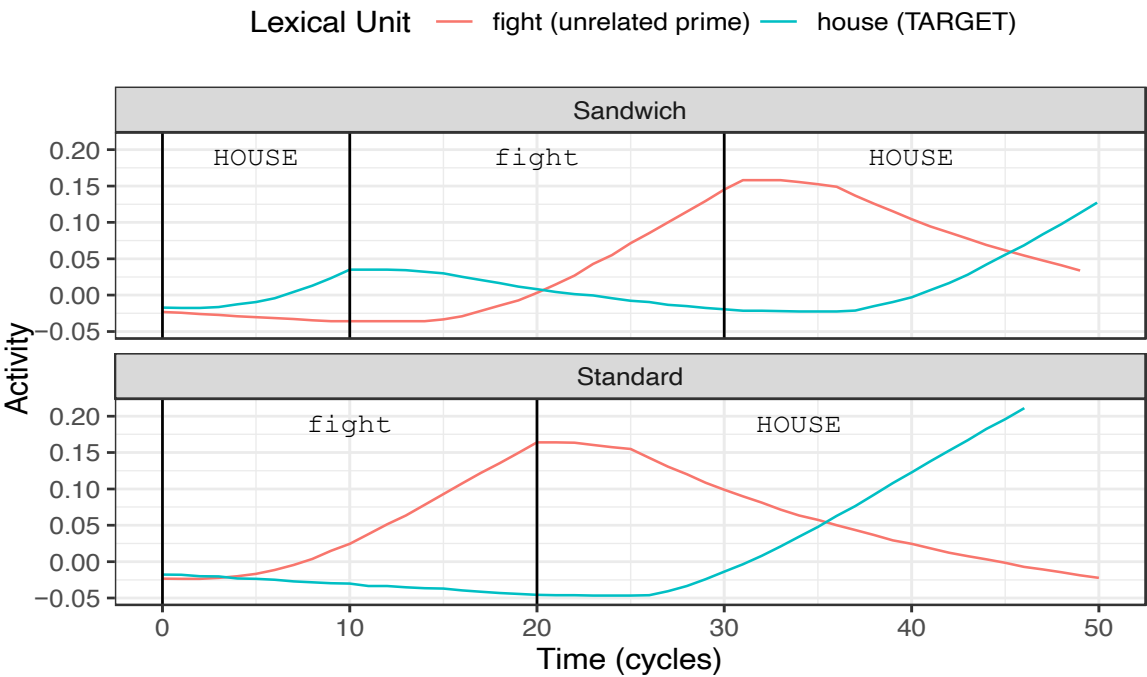
To compare the two techniques, we manipulated the procedure (standard vs. sandwich) in a within-subject design where we compared the largest and most robust form of masked priming: identity priming (i.e., identity vs. unrelated primes). Our comparison aims to shed light on the processing differences between the two techniques. To do so, we also examine the distributional features of the RT data and the relative speed of correct and error responses in the two paradigms.

For identity primes, the activation from the pre-primes in the sandwich technique would naturally help word processing: the target word is presented as a pre-prime and as a prime (e.g., compare pre-prime: `HOUSE`, prime: `house`, target: `HOUSE` in the sandwich method vs. prime: `house`; target: `HOUSE` in the standard method). Thus, word recognition times in the identity condition should be faster in the sandwich technique than in the standard technique. For the unrelated primes, the presence of the pre-prime in the sandwich technique does not generate such a straightforward prediction (e.g., pre-prime: `HOUSE`, prime: `fight`, target: `HOUSE` in the sandwich technique vs. prime: `house`, target: `HOUSE` in the standard technique). As suggested by Trifonova and Adelman (2018), in the standard technique one would expect unrelated primes to push target's activation below resting level after prime offset in interactive activation models of word recognition (e.g., SCM, Davis, 2010; see Bottom Panel of Figure 2). In contrast, because of the residual activation caused by the pre-prime in the sandwich technique, the target's activation would be higher than with the standard technique after prime offset (see Top Panel of Figure 2). Thus, one would predict that unrelated pairs would be identified faster in the sandwich than

in the standard technique—note that this mechanism would reduce the boost of the sandwich technique.

Figure 2

Example of the initial processing cycles of an unrelated trial in the sandwich technique



Note. Top panel, pre-prime: house [10 cycles], prime: fight [20 cycles], and target [until response]) and in the standard masked priming technique. Bottom panel, prime: fight [20 cycles] and target [until response]) in a leading competitive network model of visual word recognition (SCM; C. J. Davis, 2010). The thick horizontal line represents a response threshold, and the coloured lines represent the activation levels of the corresponding lexical units.

The above pattern was confirmed by simulations on the spatial coding model (Davis, 2010). We used the target words of Davis and Lupker’s (2016) Experiment 2—they were all five-letter words—paired with identity vs. unrelated primes.²⁵ For the sandwich method, we simulated a scenario with a 33-ms pre-prime followed by a 50-ms prime, whereas for the standard technique, we simulated a scenario with a 50-ms prime. We found a greater identity priming effect in the sandwich technique than in the standard technique (69 vs. 47 processing cycles). This boost was exclusively due to the identity condition (an advantage of 24 processing cycles in the sandwich technique). For the unrelated condition, the sandwich technique produced slightly faster responses than the

²⁵ The simulator is available at <http://www.pc.rhul.ac.uk/staff/c.davis/SpatialCodingModel/>

standard technique (an advantage of 2 processing cycles in the sandwich technique). Of note, in a form priming experiment, Trifonova and Adelman (2018, Experiment 2) found the opposite trend for unrelated pairs in the error rates (i.e., better performance in the standard technique). If this latter finding can be replicated, this would suggest that the way models simulate masked priming with the standard and sandwich techniques many need to be reevaluated.

To get a richer picture of the differences between the two techniques, we also aim to fully describe the distributional features of the RTs in the sandwich technique and standard methods; in the present work, we report delta plots (De Jong, Liang, & Lauber, 1994) and conditional accuracy functions (Bonnet & Dresp, 1993; Ollman, 1977). Critically, the study of RT distributions offers rich information to elucidate the time course of an effect (Balota et al., 2008; see also Heathcote et al., 1991; Ratcliff, 1979; Townsend & Ashby, 1983, for early research), because many theories of word processing make claims of the sequence and timing of mental operations, as can be seen in Figure 2. However, all previous studies with the sandwich technique only focused on the analysis of the mean RTs.

In the standard technique, masked identity priming has repeatedly shown a shift between the RT distributions of the identity and unrelated conditions (i.e., identity priming), with the magnitude of identity priming being close to that of the prime duration (e.g., Gomez & Perea, 2020; Gomez, Perea, & Ratcliff, 2013; Perea, Marcet, Lozano, & Gomez, 2018, for a similar finding). This pattern supports Forster's (1998) claim that the identity prime produces a "head start" advantage when processing the target word: the identity pair *house*-HOUSE would enjoy an advantage over the unrelated pair *fight*-HOUSE of approximately the prime duration. In contrast, when using visible primes (150 ms prime duration), identity priming elicits a shift in the RT distributions and a change in their shape as it is more skewed for the unrelated condition. This pattern is consistent with the idea that unmasked visible primes affect core decision processes (see Gomez et al., 2013). Here, we tested whether masked identity priming effects in the sandwich technique also reflect a shift in the RT distributions (i.e., a head-start advantage, as in the standard masked priming technique) or

whether they are larger in the higher quantiles (i.e., an effect also in the decision process from the identity primes, based on the evidence accumulation models). The former outcome would support the view that the two techniques tap the same core processes, whereas the latter result would imply that the sandwich technique taps core processes other than a “head-start” advantage.

In sum, to provide the first full description of the sandwich technique, we designed a lexical decision experiment that examined the similarities and differences in masked identity priming with the standard and sandwich techniques. The findings of this study are also relevant to inform us of whether the sandwich technique taps similar encoding processes (i.e., a head-start advantage) as the standard priming technique or whether the two tasks are more different than originally thought. Keep in mind that, in the sandwich technique, the pre-prime presentation adds complexity to the information processing; that is, the system must handle with three strings in brief succession (see Forster, 2009; Trifonova & Adelman, 2018). To rephrase our research question, an apt analogy might be that the two methods can be thought of as measuring devices for cognitive processes and the present study aims to understand if these devices are qualitatively different as they detect different signals, or quantitatively different, as they may magnify the same signal to different levels.

Method

Participants

A total of 36 undergraduate students from the University of València volunteered to participate in the experiment—this ensured 2,160 observations per condition, in line with Brysbaert and Stevens’ (2008) suggestion for masked priming experiments. All participants were native speakers of Spanish with normal/corrected vision and no history of reading disorders. This study was approved by the Experimental Research Ethics Committee of the University of València, and all participants signed an informed consent form.

Materials

We selected 240 word targets of 5 letters, which were extracted from Fernández-López et al.'s (2019) Experiment 1. The average Zipf frequency was 4.15 (range 3.74-4.61) and the mean OLD20 was 1.41 (range: 1.00-2.65) in the EsPal Spanish database (Duchon et al., 2013). Each target word (e.g., *PASTA* [pasta]) was paired with an identity prime (e.g., *pasta*) or an unrelated prime (e.g., *suelo* [floor]). The pre-primes in the sandwich method were always the same as the targets (see Figure 1). To act as foils, we selected 240 orthographically legal pseudowords. These were also extracted from the Fernández-López et al.'s (2019) Experiment 1. Each target nonword (e.g., *JARSA*) was paired with an identity prime (*jarsa*) or an unrelated prime (*ruezo*). We created four lists to counterbalance the prime-target combinations across the two priming conditions and the two techniques

Procedure

Testing took place in groups of 4-6 participants in a quiet lab room. Computers equipped with DMDX software (Forster & Forster, 2003) displayed the stimuli and registered the timing/accuracy of the responses. Each trial started with a pattern mask (#####) displayed for 500 ms in the center of a computer screen. In those trials with the sandwich technique, the uppercase target stimulus was presented for 33 ms as a pre-prime, followed by a 50-ms lowercase prime stimulus which, in turn, was replaced by the uppercase target stimulus—this was presented until response or 2 seconds had elapsed. In the trials with the standard technique, the pattern mask was followed by the 50-ms prime, and then replaced by the target (see Figure 1). All stimuli were presented in 14-pt Courier New black font on a white background. Participants were instructed to decide whether the uppercase stimulus was a word or not by pressing a green button (“word”) or a red button (“nonword”) with their index fingers. Both speed and accuracy were stressed in the instructions. Instructions did not mention the existence of any briefly presented primes. Sixteen practice trials preceded the 480 experimental trials. Each participant received a random sequence of trials. The session lasted 18-22 minutes.

Results

Error responses and anticipations (RTs faster than 250 ms: 2 correct responses and 4 incorrect responses) were omitted from the latency analyses. Timeouts after 2.5 seconds were categorized as errors (there were 32 timed out responses and 1,091 error responses out of 17,280 trials). Table 1 displays the mean RTs for correct responses and the accuracy in each condition. We first present the inferential analyses using Bayesian linear mixed effects models, and then we present the exploratory data analyses.

Table 1

Mean correct response times (in ms) and accuracy (in brackets) for words and pseudowords with the Standard and Sandwich masked priming techniques

	Standard technique		Sandwich technique	
	Identity	Unrelated	Identity	Unrelated
Words	589 (96.9)	629 (93.9)	563 (96.9)	642 (91.2)
Pseudowords	704 (91.6)	717 (92.6)	695 (92.7)	709 (92.0)

Bayesian linear mixed effects analyses

The RT and accuracy data were modelled with Bayesian linear mixed-effects models in R (R Core Team, 2020) using the *brms* package (Bürkner, 2020). This package, which uses the programming language Stan as its core, allows us to fit the maximal random effect structure models allowed by the experimental design (see Barr et al., 2013). The fixed effects were relation (identity vs. unrelated; coded as -0.5 and 0.5) and procedure (sandwich vs. standard; coded as -0.5 and 0.5). As usual in masked priming experiments, we conducted separate analyses for word and nonword trials.

The fitted model was as follows: Dependent variable [RT or accuracy] ~ relation * procedure + (1+relation * procedure | subject) + (1 + relation * procedure | item)

We used the ex-Gaussian family function to model the RT data and the Bernoulli family function (1 = correct, 0 = incorrect) to model the accuracy data. For the fits of each model we employed four chains, each with 5,000 iterations (1,000 warmup + 4,000 sampling). All $\hat{R}$ values were 1.00, thus meaning that the four chains converged successfully. In the output, the Bayesian linear mixed-effects models indicate the estimate of each fixed effect, its standard error, and its 95% credible interval. For inference purposes, those intervals that did not include zero were interpreted as significant. In case of a significant interaction, we computed the simple effect tests with the *emmeans* package (Lenth et al., 2020)—this package provides the 95% highest (posterior) density (HPD) interval for each effect. Note that this method yields what is called Bayesian shrinkage, which means that posterior estimates might be shifted from the empirical effects due to the prior and the data from other conditions.

Word targets

Response times were faster for identity pairs than for unrelated pairs ($b = 72.35$, $SE = 3.78$, 95% credible interval [CrI] [64.58, 79.56]) and responses were faster in the sandwich than in the standard technique ($b = 25.65$, $SE = 2.70$, 95%CrI [20.74, 31.29]). As expected, there was clear evidence of an interaction between the two factors ($b = -32.76$, $SE = 4.09$, 95%CrI [-40.80, -24.78]). This interaction reflected that, for identity pairs, responses were substantially faster in the sandwich than in the standard technique (95% HPD [-31.29, -20.7]); in contrast, for unrelated pairs, responses were faster in the standard than in the sandwich technique (95% HPD [0.75, 13.6]).

The analyses on the accuracy data showed that participants were more accurate for identity pairs than for unrelated pairs ($b = -1.16$, $SE = 0.19$, 95%CrI [-1.52, -0.79]). While the 95% credible intervals of the other parameters included zero, accuracy was lower in the sandwich than in the standard technique ($b = 0.18$, $SE = 0.23$, 95%CrI [-0.28, 0.65]) and the size of the identity priming effect

was greater in the sandwich than in the standard technique (interaction: $b = 0.40$, $SE = 0.26$, 95%CrI [-0.11, 0.90]).

Nonword targets

Responses to nonwords were faster in the trials with the sandwich technique than in the trials with the standard technique ($b = 17.54$, $SE = 3.67$, 95%CrI [10.37, 24.72]). We also found a small advantage of identity over unrelated pairs ($b = 9.83$, $SE = 5.07$, 95%CrI [0.22, 19.93]). There were no clear signs of an interaction between the two factors ($b = -4.21$, $SE = 5.11$, 95%CrI [-14.30, 5.7]).

The analyses on the accuracy data did not show any effects (relation: $b = -0.14$, $SE = 0.16$, 95%CrI [-0.45, 0.18]; procedure: $b = -0.14$, $SE = 0.16$, 95%CrI [-0.45, 0.17]; relation*procedure: $b = 0.21$, $SE = 0.21$, 95%CrI [-0.20, 0.61]).

To summarise, the results with the Bayesian linear mixed-effects models on word trials showed that, for identity primes, the presentation of the target word in the sandwich technique yielded faster responses than the standard technique (HOUSE-house-HOUSE < house-HOUSE). In contrast, the sandwich technique slowed down word responses for unrelated primes (fight-HOUSE < HOUSE-fight-HOUSE). Thus, the boost in the sandwich technique is caused by faster responses in the identity condition and slower responses in the unrelated condition. Regarding the nonword trials, we found faster responses in the sandwich technique, accompanied by a small identity priming effect.

Exploratory Data Analyses

To further examine the nature of the priming effects underlying the sandwich and standard techniques, we compared the empirical RT distributions for both techniques. We focused on delta plots and conditional accuracy functions, as they describe how the latency distributions differ. Note that these methods do not have established inferential properties, so in our discussion we focus only on large effects.

Delta plots

Delta plots are frequently used to display how a latency effect (e.g., identity priming, or task) evolves across time (see De Jong et al., 1994; Ridderinkhof et al., 2004). These plots are constructed as follows:

1. We obtain the RTs at the desired quantiles (here we use .1, .2, .3, .4, .5,9) for every participant for each of the conditions to be compared, only for correct responses.
2. We average the RTs obtained in Step 1 across participants (these averaged quantiles are also called vincentiles in the RT literature; see Ratcliff, 1979).
3. For each of the quantiles in the vincentiles, (a) we find the average between the two conditions, and (b) then compute the differences (i.e., the delta) preserving the sign.
4. The last step is to plot the points (as many points as there are quantiles), with the averages (Step 3a) on the x-axis, and the delta (Step 3b) on the y-axis.

As is evident from Step 3b, delta plots are based on residual quantiles (i.e., RTs at quantiles in condition A minus condition B), and hence can provide us with some indication of the temporal dynamics of an effect when interpreted within the lens of process models. For example, a flat line at around $y = 50$ ms would mean that there is a 50-ms shift in the RT distributions—this could be interpreted as faster encoding times in evidence accumulation models (see Gomez et al., 2013, for an example of such interpretation in the standard masked priming technique). Alternatively, an ascending function would indicate that the effect grows for slower responses, and it could be interpreted as a difference in the rate of evidence accumulation.

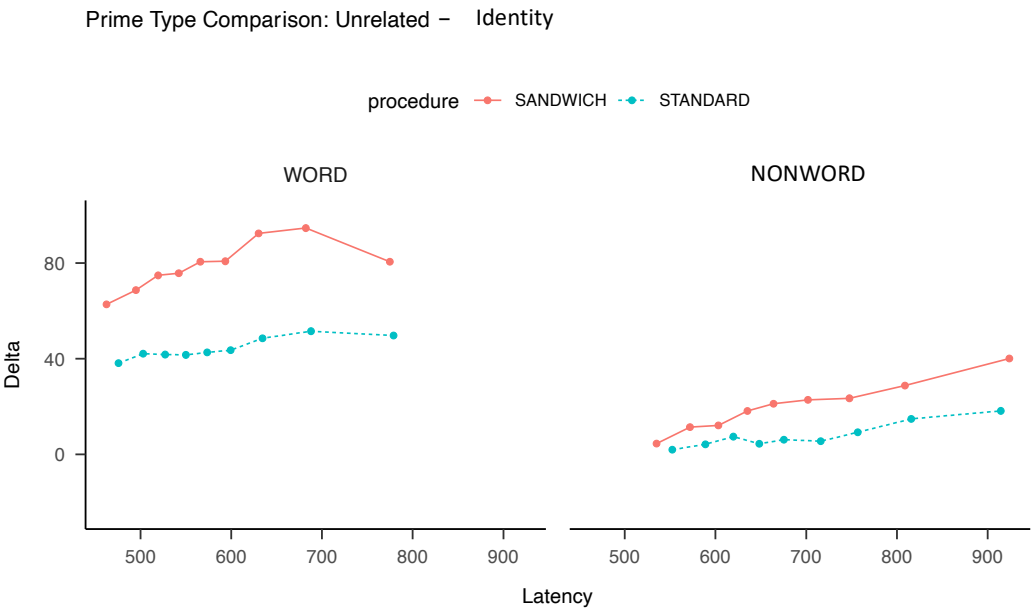
As just indicated, in delta plots, one needs to contrast two conditions; in our analysis of the two priming methods, it makes sense to perform two separate comparisons that are presented in two different delta plots: (a) the effect of prime type in the standard and sandwich techniques (across condition comparison:

unrelated vs. identity; Figure 3), and (b) the effect of procedure for identity and unrelated pairs (within-condition comparison: standard vs. sandwich; Figure 4).

Effect of prime type (Figure 3). For word trials, when calculating priming effects as RT for unrelated primes minus RT for identity primes, we found a relatively flat line for the identity priming effect in the standard technique (all points lie between 38 and 51 ms), thus replicating earlier research. In contrast, in the sandwich technique, the identity priming effect was not only greater than in the standard technique (as attested by the higher values in the y-axis), but it progressively increased from the .1 quantile (63 ms) to the .7 quantile (95 ms)—it decreased slightly in the .9 quantile, though. For the nonword trials, we found a small identity priming effect that slowly grew across quantiles, with a steeper slope in the sandwich technique.

Figure 3

Delta plot comparing unrelated primes to identity primes



Note. The points represent the difference RT unrelated – RT identity at the .1, .2, .3, .4, .5, .6, .7, .8, and .9 quantiles

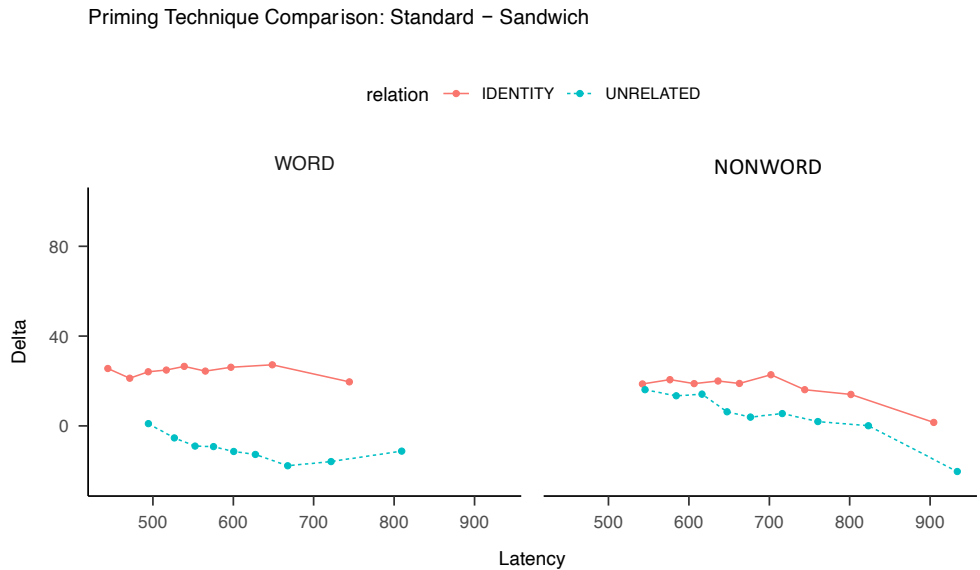
Effect of technique (Figure 4). For word trials, when calculating the effect of technique as *standard* minus *sandwich*, we found an approximate flat line for the identity primes (around 25 ± 5 ms for all quantiles). This pattern indicates that responses were overall *faster* in the sandwich than in the standard technique. In

contrast, we found a line that lies on the negative numbers for the unrelated primes (i.e., responses were *slower* in the sandwich than in the standard technique)—note that this difference effect grew as quantiles increased from .1 quantile (-.5 ms) to .8 quantile (-.16 ms; i.e., a negative slope). We found functions with negative slopes for the nonword targets, but with points above zero for the most part. This pattern reflected that responses were faster in the sandwich technique, but this advantage decreases (or even reversed, in the unrelated condition) for the slower responses.

Note that delta plots highlight the dynamics of effects across different points of the latency distributions. However, they do not convey information about accuracy as they are only built with correct responses; hence, we resort to conditional accuracy functions to explore accuracy issues as explained below.

Figure 4

Delta plot comparing the standard and sandwich techniques



Note. The points represent the difference RT standard – RT sandwich at the .1, .2, .3, .4, .5, .6, .7, .8, and .9 quantiles

Conditional Accuracy Functions

Conditional accuracy functions (Bonnet & Dresch, 1993; Ollman, 1977) allow us to visualize how accuracy evolves as latencies increase. These plots were generated as follows:

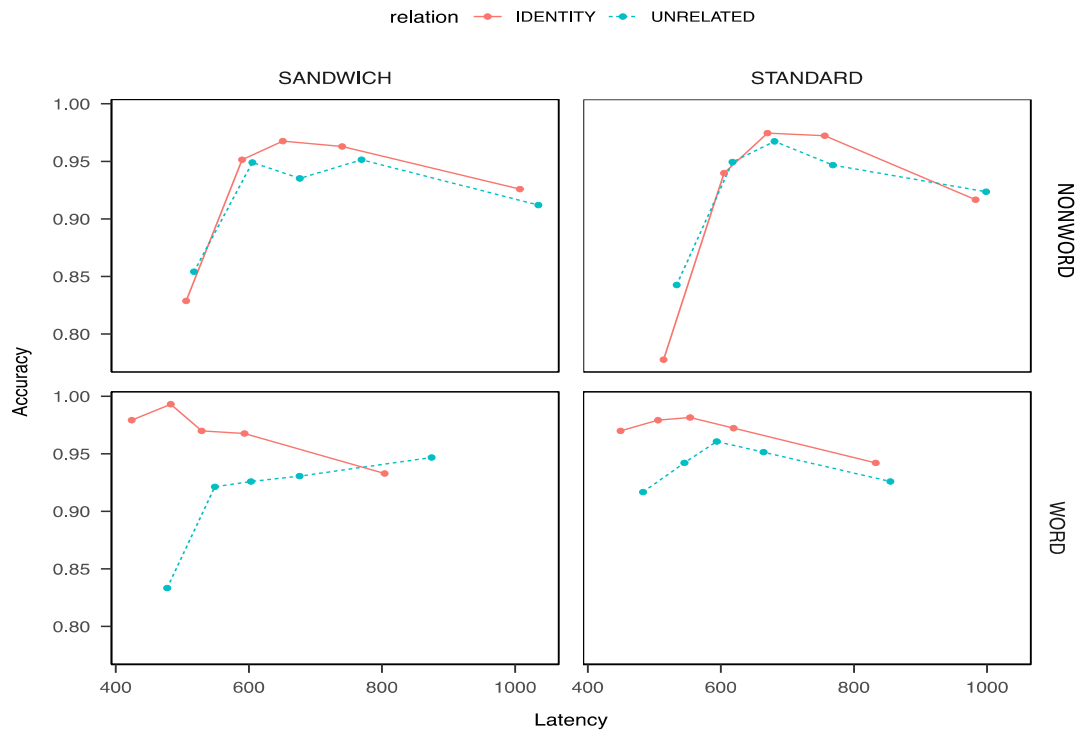
1. For each type of item and each participant, we found the RTs at quantiles .1, .3, .5, .7, and .9; to do so, we included both correct and error responses.
2. We averaged the RTs and accuracies for each bin across all participants, and then calculate (a) the average RT for each of the bins, and (b) the accuracy within each bin.
3. The last step is to plot the points (as many points as there are quantiles), with the averages (Step 3a) on the *x*-axis, and the conditional accuracies (Step 3b) on the *y*-axis.

As is evident from Step 3b, conditional accuracy functions can indicate if the accuracy in the responses for a given condition varies across the latency of the responses to such condition. For example, a flat line at around $y = .90$ would mean that the accuracy for that condition (90%) is equal for fast and for slow responses. In addition, an ascending function would indicate that the fast responses tend to be less accurate, and a descending function would indicate that slow responses tend to be less accurate.

The panels in Figure 5 show the conditional accuracy for all the conditions in our experiment. Two patterns are worth highlighting: for words, the fastest 20% of responses in the unrelated condition in the sandwich method tend to be less accurate than slower responses, which is reversed for the identity condition. For nonwords, the conditional accuracy functions are remarkably similar in the two techniques: fastest responses tend to be less accurate for the identity primes than for the unrelated primes, but this pattern reverses for the slower responses. This is consistent with the idea that familiarity influences lexical decision (e.g., Masson & Bodner; cf compound-cues), particularly for fast responses, which could be sensitive to the orthographic match of the prime and target. A match could bias for “word” responses, and fast responses might be most affected by this bias.

Figure 5

Conditional accuracy functions



Note. The points represent the accuracy and average RT of the responses within equal sized bins (20% of responses per bin)

Discussion

An increasing number of masked priming researchers now opt for the sandwich technique due to its promise of greater effect sizes than the standard technique. However, the mechanisms underlying this boost are still not well understood. To fully characterize the sandwich technique, we conducted a lexical decision experiment comparing the sandwich and standard techniques in a within-subject design. We focused on masked identity priming effects for three reasons: (a) theoretical models offer straightforward predictions (Forster, 1998); (b) it has already been modelled with the standard masked priming technique (Gomez et al., 2013); and (c) the magnitude of the effects is much larger than with other types of prime-target relationships (i.e., we may capture more easily potential interactions). To elucidate whether the differences between the standard and sandwich methods were merely quantitative or also qualitative, we examined the

locus of the boost in the sandwich technique and the time course of the effects for both techniques via delta plots and conditional accuracy functions.

For word trials, we found the expected boost in the sandwich technique (a 79-ms priming effect with the sandwich technique vs. a 40-ms priming effect with the standard method; i.e. a 39-ms boost), thus replicating earlier research (see Lupker & Davis, 2009, for the first demonstration). Critically, the use of a within-subject design allowed us to discover that while the major component of the 39-ms boost in the magnitude of priming effects in the sandwich technique was due to the faster responses to identity pairs (26-ms: 563-ms vs. 589-ms in the sandwich and standard techniques, respectively), there was also a component due to the slower responses to unrelated pairs (13-ms: 642-ms vs. 629-ms, in the sandwich and standard techniques, respectively)—note that Trifonova and Adelman (2018) also found an effect in this direction for unrelated pairs in the accuracy data. As indicated in the Introduction, Davis' (2010) spatial coding model would have predicted not only substantially faster responding in the identity condition in the sandwich than in the standard technique, but also slightly faster responding in the unrelated condition. Thus, the effect of the unrelated priming conditions predicted by the spatial coding model does not match the observed data. We found that the responses were slower in the sandwich technique than in the standard technique after an unrelated prime. The discrepancy in the unrelated trials suggests that the pre-prime in the sandwich technique may influence target processing in a way that cannot be captured in the present implementations of masked priming in interactive activation models like the spatial coding model.

In addition, the exploratory data analyses (both delta plots and conditional accuracy functions) offer valuable information on the underlying processing differences between the standard and the sandwich techniques. For word targets, the delta plot on the standard technique showed, as expected, that masked identity priming reflects essentially a shift in the RT distributions (see Figure 3; see also Gomez et al., 2013; Gomez & Perea, 2020). This pattern is entirely consistent with the commonly shared assumption that masked identity primes have a “head start” advantage during an early encoding stage (i.e., a “savings”

effect; Forster, 1998; see also Gomez et al., 2013). In contrast, the delta plot with the sandwich technique showed a different picture: the magnitude of masked identity priming grows stronger in the higher quantiles (see Figure 3). At first glance, one might interpret this last pattern as being due to some involvement of decision processes for identity primes in the sandwich technique. However, this interpretation misses the critical information provided by within-condition comparisons (see Jacobs et al., 1995). The advantage of the identity primes in the sandwich technique is approximately constant along quantiles (see left panel of Figure 4), thus reinforcing the idea that the identity primes are processed similarly in the two techniques (i.e., an encoding effect; Forster, 1998). Instead, the largest differences across tasks occur in the unrelated priming condition, where there is a costlier processing in the sandwich technique as latency increases. Of note, the delta plot in the left panel of Figure 4 also allows us to visualize that the difference between the sandwich and the standard method manifests itself in both identity and unrelated primes: the sandwich technique yields faster responses for the identity primes and slower responses for the unrelated primes.

The Conditional accuracy functions shown in Figure 5 provide further clues on the differences between the two techniques. The fast responses in the sandwich method for the unrelated condition are much less precise than the fast responses for any other condition in the experiment. This pattern suggests that the pre-prime interacts with the prime in ways that affect the early moments of processing of the target. One possible explanation for the divergent pattern in the unrelated primes across technique is the following. Because of the briefness of the pre-prime, its processing may overlap with the prime, creating a jumbled orthographic code that would impair the first moments of target processing—note that this interpretation has some resemblance with the legibility effect described by Davis and Forster (1994). For instance, for `HOUSE-fight-HOUSE` (sandwich technique), when participants initially encounter `HOUSE-fight`, they might perceive a fused orthographic code that might initially hinder target recognition, occasionally leading to fast error responses. This early difficulty would be later

overwhelmed by lexical activation.²⁶ Importantly, this is a post hoc explanation that can be interpreted in different ways. What do we mean by a “jumbled orthographic code”? Two implementations of this idea come to mind: the first one is that the orthographic features of both the pre-prime and the target are partially available, and hence there might be at least some benefit for the sandwich technique over the standard technique; unfortunately, this idea does not seem to be supported by the data because RTs for the sandwich technique are slightly longer than for the standard technique. Another way to think about this “jumbling” is that the orthographic codes mutually inhibit each, thus diminishing lexical activation for both entries. This question, however, is well beyond the scope of the present article. Nevertheless, it raises interesting issues around how information from the strings integrates versus contrast in the type of fast presentation displays in this type of paradigms. These questions will be central in future analyses of the sandwich technique, which was developed to explore rather fine-grained orthographic processes that yield much smaller effects than identity versus unrelated primes.

In sum, the present experiment has both theoretical and practical implications. On the theoretical side, the shift in the RT distributions for identity pairs in the sandwich and standard technique strongly suggests that the extra time from the pre-prime is not altering the underlying processes beyond encoding. Notably, the story is different from unrelated primes: the pre-prime in the sandwich technique hinders the processing of unrelated prime-target pairs. Further research with a better temporal resolution (i.e., ERPs) is necessary to uncover these differences—note that the Ktori et al. (2012) ERP sandwich experiment did not include a block with the standard technique. On the applied side, one basic question is shall we continue using the sandwich technique? The answer from the present experiment is “yes”. Even in the largest manipulation in priming experiments (i.e., identity primes), we found a similar pattern in the two

²⁶ While the main focus of this work is the words, no complete account of the tasks can exclude a description of the responses to pseudowords. In our data, the conditional accuracy function and the delta plots are remarkably similar, with effects augmented in the sandwich condition.

techniques for identity pairs—except for the expected processing advantage in the sandwich technique due to the pre-prime. At the same time, researchers should be cautious in interpreting some of the across condition priming effects in subtle manipulations in the sandwich technique, as these effects may derive not from facilitation in the related condition but rather from some inhibition in the unrelated condition and might consider using both the identity and unrelated conditions as controls when exploring subtle manipulations (e.g., does `corn` prime `ACORN`?) Perhaps an analogy from the physical sciences is apt. We explored if the sandwich technique is a more powerful microscope than the standard masked priming technique. Our results indicate that a better way to think about it is as a different catalyst that triggers related but slightly different processes.

Conclusions

This section is composed of two studies that aimed to examine in depth the mechanisms underlying the masked priming technique and the sandwich technique. In the first study, *Can response congruency effects be obtained in masked priming lexical decision?* (Fernández-López et al., 2019), we examined whether participants applied the same instructions to the primes and targets in the masked priming lexical decision task. To that end, we focused on the scrutiny of the congruency effect (i.e., the difference in response times between a congruent trial (bright-WINDOW: “word”; geuzel-POLCIR: “nonword”) and an incongruent trial (geuzel-WINDOW: “word”; bright-POLCIR: “nonword”).

We designed three experiments to examine the influence of the lexical status and wordlikeness of unrelated primes in lexical decision. The primes could be a high-frequency word (radio), a low-frequency word (beret), a pseudoword (geuze), or a consonant string (xfgtp). To obtain a comprehensive picture, we developed three scenarios that varied in the difficulty of the word/nonword discrimination: orthographically legal nonwords matched in subsyllabic elements with the target words (Experiment 1; e.g., WINDOW vs POLCIR), orthographically illegal nonword foils (Experiment 2; e.g., WINDOW vs SYKDD), and orthographically legal nonwords with no neighbors (Experiment 3; e.g., WINDOW vs LEDUL). We found that when using orthographically legal nonwords (POLCIR, LEDUL), response times to words and nonwords did not differ as a function of the wordlikeness of the prime stimuli. This suggests that participants were unable to establish the lexical status of the masked prime or even to determine the likelihood of its being a word. Nevertheless, we found a different scenario when using illegal nonwords as foils (WINDOW vs SYKDD): word response times were longer when the masked prime was a consonant string than when the masked prime was an orthographically legal nonword (xfgtp-WINDOW > geuzel-WINDOW). Moreover, nonword responses were faster when preceded by a consonant-string prime than when preceded by an orthographically legal prime (xfgtp-SKYDD < geuzel-SKYDD), but only in the initial quantiles of RT

distributions. Nevertheless, one might argue that the inclusion of illegal foils in experiment 2 switched the nature of the task from a lexical decision to an orthographic legality decision task.

These findings have important implications for leading models of word recognition (Bayesian Reader model; Norris & Kinoshita, 2008; interactive activation models; Rumelhart, 1981; Grainger & Jacobs, 1996; Coltheart et al., 2001; Davis, 2010). The Bayesian Reader (BR) model predicts similar response times for words preceded by an unrelated prime regardless of its wordlikeness in lexical decision. Consistent with the BR model, all unrelated priming conditions behaved similarly in standard word recognition scenarios. However, when the word/nonword discrimination was much easier (Experiment 2; WINDOW vs SYKDD), we found an effect from the unrelated primes. ($xfgtp-WINDOW > geuzel-WINDOW$; $xfgtp-SKYDD < geuzel-SKYDD$). Nonetheless, since the task in Experiment 2 can be regarded as an orthographic legality decision, the BR model can predict slower “yes” responses when the prime is composed by a consonant string than when the prime is an orthographically legal stimulus, and vice versa. As all the nonword foils violated the orthographic constraints, anything legal must be a word. Regarding the interactive activation models, the lack of influence of the wordlikeness of the prime on the target processing when comparing words and legal foils (Experiment 1 and 3) offers support to the selective lateral inhibition hypothesis (Davis & Lupker, 2006). However, the nature of the task in Experiment 2 (i.e., orthographic legality decision) would go beyond the scope of the family of interactive activation models. Thus, our findings offered empirical support to the BR account and the lateral inhibition hypothesis of masked priming in standard lexical decision experiments. Importantly, they also demonstrate how the nature of the decision required by the task modulates masked priming effects.

Notably, the masked priming effects can also be modulated by modifying the technique. With the purpose of magnifying the priming effects, Lupker and Davis (2009) designed the sandwich masked priming paradigm, in which the forward mask is followed by an initial masked prime identical to the target (the so-called pre-prime) for 33 ms, followed by the prime of interest, and the

subsequent target (####-WINDOW-geuzel-WINDOW). Whereas the sandwich priming technique has been employed in a huge number of experiments due to the greater effect sizes, the mechanisms underlying this boost were still not well understood. In the second study of the present section, *Unveiling the boost in the sandwich priming technique* (Fernández-López et al., 2021), we aimed to fully characterize the sandwich technique. To that end, we conducted an identity priming lexical decision experiment comparing the sandwich and standard techniques in a within-subject design. Moreover, to test the nature of the differences between the standard and sandwich methods (i.e., whether they were merely quantitative or also qualitative), we examined the locus of the boost in the sandwich technique and the time course of the effects for both techniques via delta plots and conditional accuracy functions. Results showed that the locus of the boost in the sandwich technique was two-fold: faster responses in the identity condition (via a shift in the RT distributions) and slower responses in the unrelated condition. The shift in the RT distributions for identity pairs in the sandwich and standard technique strongly suggests that the extra time from the pre-prime is not altering the underlying processes beyond encoding. Nevertheless, for unrelated primes, the pre-prime in the sandwich technique hampered the processing of unrelated prime-target pairs. This finding warns researchers to be cautious in interpreting some of the across condition priming effects in subtle manipulations in the sandwich technique, as these effects may derive not from facilitation in the related condition but rather from inhibition in the unrelated condition.

In summary, the two studies presented here offer a valuable contribution to the methodological aspects of word recognition research. Three highlights can be drawn from them: (i) when choosing the baseline, it does not matter the wordlikeness of the unrelated primes (as long as the task employs orthographically legal foils); (ii) the nature of the decision required by the task modulates masked priming effects; (iii) when employing the sandwich technique, researchers should use both the identity and unrelated conditions as controls when exploring subtle manipulations.

Concluding Remarks

If one thing strikes us as reading researchers, it is how the brain has developed specialized cortical mechanisms attuned to the recognition of written words. Letters are special visual objects, and as I have indicated as a mantra in this dissertation, they can be considered tools that modify the brain machinery that sustains reading. But, how does this occur? The backbone of the present dissertation is the question of how literacy shapes the cognitive processes that underlie the encoding of visual perception to meaning (i.e., orthographic processing).

Orthographic processing is the mechanism responsible for encoding the identity and position of letters within words, allowing us not only to distinguish `silent` from `listen` but also to access the lexical entry of `rabbit` independently of whether the word is written as *rabbít*, *rabbit*, *rabbit*, or *rabbít*. Notably, recognizing these words is necessarily mediated by the processing involving the letters that make up the word, at least in languages that use alphabetic orthography. Indeed, it is generally assumed that letters are the key elements for orthographic processing and that it is the mechanism used to code for the positions of these letters that critically determines the nature of the orthographic code.

In the first section of the present dissertation, *The emergence of the orthographic processing*, I aimed to further our understanding of how the encoding of the position and order of letters emerges with literacy. We found evidence demonstrating that even though our brain is not phylogenetically prepared for reading, the specialization to letters emerges very early, with the mere visual exposure. Specifically, the mechanisms of visual object processing are attuned to guide the functional reading process and develop some rudimentary specialization to letters before the actual emergence of orthographic processing.

On the applied side, these outcomes highlight the importance of the assessment of children's analysis of visual input, which is usually overlooked, in order to detect early deficiencies and prevent eventual reading difficulties. Moreover, the learning to read instruction should not be based solely on letter knowledge and phonological decoding. We must keep in mind that reading is built on a basic cognitive foundation. And that foundation can be trained from very early in development to smooth the learning-to-read process. It is advisable to train in visual tasks with letter stimuli, like the same-different task, to facilitate the encoding of letter order and position. In addition, these findings suggest that children may benefit to a relatively early exposition to words. This would help create a rudimentary layer of frequent letter co-occurrences that may facilitate the subsequent encoding of words.

As learning to read becomes more automatic, and with the final goal of efficient word recognition, the reading brain becomes more flexible to letter order (and I bet you have read order) and more tolerant to the differences in letter shape. The second section of this thesis, *The resilience of letter detectors*, focuses on the tolerance of letter detectors to letter distortion. We focused on letter rotations, as leading models of word recognition have explicit predictions about how the processing system copes with them. We found that the letter detectors responsible for converting the visual input into abstract representations are hindered by rotations since very early in processing. Nevertheless, they can efficiently overcome the constraints imposed by letter rotations up to 45°. After this boundary, individuals need to engage extra top-down (lexical) feedback, which results in larger identification times. This pattern has implications at a theoretical level. The SERIOL model of word recognition (Whitney, 2001, 2002) assumes that "input levels to letter units are reduced for rotated stimuli" (p. 119, Whitney, 2002), increasing gradually the time required to encode a letter string. Instead, the Local Combination Detectors model (Dehaene et al., 2005) offers an all-or-none account of how letter rotation affects letter and word processing. It postulates a negligible cost for letter rotation until a boundary of 40°-45°, beyond which letter detectors would not efficiently cope with rotations. Based on the results obtained in the different experiments of this section, the most feasible

alternative is the combination of the assumptions of both models. There is a gradual cost of letter rotation until a limit, around 45°, beyond which the letter detectors' resilience decreases, becoming dramatic at steeper angles.

Over these years as a graduate student, I have touched on research topics other than orthographic processing. As a result, there are the studies that compose the third section of the present thesis, *Methodological contributions to the visual word recognition research*, which examined the cognitive underpinnings of masked priming and one of its variants: sandwich priming. The results obtained in this section have important methodological implications, especially when choosing the appropriate baseline in identity priming experiments. Firstly, we have demonstrated that participants cannot establish the lexical status of the masked prime or even determine the likelihood of its being a word in a standard lexical decision task. Thus, it does not matter the wordlikeness of the unrelated prime in identity masked priming lexical decision. Secondly, researchers should use both the identity and unrelated conditions as controls when exploring subtle manipulations with the sandwich masked priming technique. The reason is that the effects could derive not from the facilitation in the related condition but rather from inhibition in the unrelated condition.

Overall, the present dissertation develops a robust understanding of orthographic processing in particular and word recognition in general, contributing to the field in several ways, with important empirical, theoretical, methodological, and practical implications. Most importantly, probably, is the thorough knowledge I have acquired from its completion, not only about the cognitive correlates of reading but also about science. Science is rigid but flexible, exact but refutable, and incredibly reflexive. Science is not competition (as much as they would have us believe the opposite), but cooperation. We all run this race with the same ultimate goal, that of broadening and strengthening our knowledge. For this reason, all the materials, data files, and data analyses are shared in the Open Science Framework, where also the published manuscripts are available. Since its inception, letters have allowed the democratic transmission of knowledge. Let's follow that line.

Todos somos nada; sin las palabras, dime, ¿qué nos queda?

Xoel López, *Tierra*

References²⁷

Resumen

- Benyhe, A., & Csibri, P. (2021). Can rotated words be processed automatically? Evidence from rotated repetition priming. *Memory & Cognition*, 1–9. doi:10.3758/s13421-021-01147-4
- Chetail, F. (2017). What do we do with what we learn? Statistical learning of orthographic regularities impacts written word processing. *Cognition*, 163, 103–120. doi:10.1016/j.cognition.2017.02.015
- Coltheart, M., Rastle, K., Perry, C., Langdon, R., & Ziegler, J. (2001). DRC: a dual route cascaded model of visual word recognition and reading aloud. *Psychological Review*, 108, 204–256. doi:10.1037/0033-295X.108.1.204
- Dandurand, F., Grainger, J., & Dufau, S. (2010). Learning location-invariant orthographic representations for printed words. *Connection Science*, 22, 25–42. doi: 10.1080/09540090903085768
- Davis, C. J. (2010). The spatial coding model of visual word identification. *Psychological Review*, 117, 713–758. doi:10.1037/a0019738
- Dehaene, S., Cohen, L., Sigman, M., & Vinckier, F. (2005). The neural code for written words: A proposal. *Trends in Cognitive Sciences*, 9, 335–341. doi:10.1016/j.tics.2005.05.004
- Duñabeitia, J. A., Dimitropoulou, M., Grainger, J., Hernández, J. A., & Carreiras, M. (2012). Differential sensitivity of letters, numbers, and symbols to character transpositions. *Journal of Cognitive Neuroscience*, 24, 1610–1624. doi:10.1162/jocn_a_00180

²⁷ References are listed by section

- Duñabeitia, J. A., Lallier, M., Paz-Alonso, P. M., & Carreiras, M. (2015). The impact of literacy on position uncertainty. *Psychological Science*, 26, 548–550. doi:10.1177/0956797615569890
- Duñabeitia, J. A., Orihuela, K., & Carreiras, M. (2014). Orthographic coding in illiterates and literates. *Psychological Science*, 25, 1275–1280. doi:10.1177/0956797614531026
- Fernández-López, M., Davis, C. J., Perea, M., Marcet, A., & Gómez, P. (in press). Unveiling the boost in the sandwich priming technique. *Quarterly Journal of Experimental Psychology*. doi:10.1177/17470218211055097
- Fernández-López, M., Gómez, P., & Perea, M. (2021). Which factors modulate letter position coding in pre-literate children? *Frontiers in Psychology*, 12, 708274. doi:10.3389/fpsyg.2021.708274
- Fernández-López, M., Marcet, A., & Perea, M. (2019). Can response congruency effects be obtained in masked priming lexical decision? *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 45, 1683–1702. doi:10.1037/xlm0000666
- Fernández-López, M., Marcet, A., & Perea, M. (2021). Does orthographic processing emerge rapidly after learning a new script? *British Journal of Psychology*, 112, 52–91. doi:10.1111/BJOP.12469
- Fernández-López, M., Mirault, J., Grainger, J., & Perea, M. (2021). How resilient is reading to letter rotations? *A parafoveal preview investigation. Journal of Experimental Psychology: Learning, Memory, and Cognition*, 47, 2029–2042. doi:10.1037/xlm0000999
- Forster, K. I., & Davis, C. (1984). Repetition priming and frequency attenuation in lexical access. *Journal of experimental psychology: Learning, Memory, and Cognition*, 10, 680–698. doi:10.1037/0278-7393.10.4.680
- Gomez, P., & Perea, M. (2020). Masked identity priming reflects an encoding advantage in developing readers. *Journal of Experimental Child Psychology*, 199, 104911. doi:10.1016/j.jecp.2020.104911

- Gómez, P., Marcet, A., & Perea, M. (2021). Are better young readers more likely to confuse their mother with their mother? *Quarterly Journal of Experimental Psychology*, 74, 1542–1552. doi:10.1177/17470218211012960
- Gomez, P., Perea, M., & Ratcliff, R. (2013). A diffusion model account of masked versus unmasked priming: Are they qualitatively different? *Journal of Experimental Psychology: Human Perception and Performance*, 39, 1731–1740. doi:10.1037/a0032333
- Grainger, J. (2008). Cracking the orthographic code: An introduction. *Language and Cognitive Processes*, 23, 1–35. doi:10.1080/01690960701578013
- Grainger, J. (2018). Orthographic processing: A ‘mid-level’ vision of reading. *Quarterly Journal of Experimental Psychology*, 71, 335–359. doi:10.1080/17470218.2017.1314515
- Grainger, J., & Dufau, S. (2012). The front-end of visual word recognition, in J.S. Adelman (Ed). *Visual word recognition, Vol. 1* (pp. 159–184). Hove: Psychology Press.
- Grainger, J., & Jacobs, A. M. (1996). Orthographic processing in visual word recognition: a multiple read-out model. *Psychological Review*, 103, 518–565. doi:10.1037/0033-295X.103.3.518
- Huey, E. B. (1908). *The psychology and pedagogy of reading*. New York: McMillan. Republished in 1968. MIT Press.
- Jackson, N., & Coltheart, M. (2001). Routes to reading success and failure. Hove, UK: Psychology Press.
- Loth, S., & Davis, C. J. (2010). Testing computational accounts of response congruency in lexical decision. In E. J., Davelaar (Ed)., *Connectionist models of neurocognition and emergent behavior: From theory to applications* (pp. 173–189). Singapore: World Scientific Publishing. doi:10.1142/9789814340359_0012
- Lupker, S. J., & Davis, C. J. (2009). Sandwich priming: A method for overcoming the limitations of masked priming by reducing lexical competitor

- effects. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 35, 618–639. doi:10.1037/a0015278
- Marcet, A., Perea, M., Baciero, A., & Gomez, P. (2019). Can letter position encoding be modified by visual perceptual elements? *Quarterly Journal of Experimental Psychology*, 72, 1344–1353. doi:10.1177/1747021818789876
- Massol, S., Duñabeitia, J. A., Carreiras, M., & Grainger, J. (2013). Evidence for Letter-Specific Position Coding Mechanisms. *PLOS ONE*, 8, e68460. doi:10.1371/journal.pone.0068460
- Norris, D., & Kinoshita, S. (2008). Perception as evidence accumulation and Bayesian inference: Insights from masked priming. *Journal of Experimental Psychology: General*, 137, 434–455. doi:10.1037/a0012799
- Pagán, A., Blythe, H. I., & Liversedge, S. P. (2021). The influence of children's reading ability on initial letter position encoding during a reading-like task. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 47, 1186–1203. doi:10.1037/xlm00 00989
- Perea, M., Jiménez, M., & Gomez, P. (2016). Does location uncertainty in letter position coding emerge because of literacy training? *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 42, 996–1001. doi:10.1037/xlm0000208
- Perea, M., Marcet, A., & Fernández-López, M. (2018). Does letter rotation slow down orthographic processing in word recognition? *Psychonomic Bulletin and Review*, 25, 2295–2300. doi:10.3758/s13423-017-1428-z
- Pollatsek, A., & Treiman, R. (Eds.). (2015). *The Oxford handbook of reading*. Oxford University Press
- Rayner, K. (1975). Parafoveal identification during a fixation in reading. *Acta Psychologica*, 39, 271–281. doi:10.1016/0001-6918(75)90011-6
- Rubinstein, H. Garfield, L., & Millikan, J. A. (1970). Homographic entries in the internal lexicon. *Journal of Verbal Learning and Verbal Behavior*, 9, 487–494. doi:10.1016/S0022-5371(70)80091-3

- Rubinstein, H., Lewis, S. S. & Rubinstein, M. A. (1971a). Homographic entries in the internal lexicon: Effects of systematicity and relative frequency of meanings. *Journal of Verbal Learning and Verbal Behavior*, 10, 57–62. doi:10.1016/S0022-5371(71)80094-4
- Rubinstein, H., Lewis, S. S. & Rubinstein, M. A. (1971b). Evidence for phonemic recoding in visual word recognition. *Journal of Verbal Learning and Verbal Behavior*, 10, 645–657. doi:10.1016/S0022-5371(71)80071-3
- Rumelhart, D. E. (1985). Toward an interactive model of reading. In S. Dornic (Ed.), *Attention and performance VI* (pp. 573–603). Erlbaum.
- Sellés, P., Martínez, T., Vidal-Abarca, E., & Gilabert, R. (2008). *BIL 3–6. Batería de Inicio a la Lectura* [Battery to Assess the Abilities Related with Early Reading Acquisition]. Madrid, Spain: Instituto de Ciencias de la Educación.
- Trifonova, I. V., & Adelman, J. S. (2018). The sandwich priming paradigm does not reduce lexical competitor effects. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 44, 1743–1764. doi:10.1037/xlm0000542
- Tydgat, I., & Grainger, J. (2009). Serial position effects in the identification of letters, digits, and symbols. *Journal of Experimental Psychology: Human Perception and Performance*, 35, 480–498. doi:10.1037/a0013027
- Vergara-Martínez, M., Gutierrez-Sigut, E., Perea, M., Gil-López, C., & Carreiras, M. (2021). The time course of processing handwritten words: An ERP investigation. *Neuropsychologia*, 159, 107924. doi:10.1016/j.neuropsychologia.2021.107924
- Whitney, C. (2001). How the brain encodes the order of letters in a printed word: The SERIOL model and selective literature review. *Psychonomic Bulletin & Review*, 8, 221–243. doi:10.3758/BF03196158
- Whitney, C. (2002). An explanation of the length effect for rotated words. *Cognitive Systems Research*, 3, 113–119. doi:10.1016/S1389-0417(01)00050-X

Yang, H., & Lupker, S. J. (2019). Does letter rotation decrease transposed letter priming effects? *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 45, 2309–2318 doi:10.1037/xlm0000697

Introduction

Dehaene, S., Cohen, L., Sigman, M., & Vinckier, F. (2005). The neural code for written words: A proposal. *Trends in Cognitive Sciences*, 9, 335–341. doi:10.1016/j.tics.2005.05.004

Fehrenbacher, D. E. (Ed.). (1989). *Lincoln: Speeches and writings*. New York, NY: Library of America.

The emergence of the orthographic processing

Andrews, S., & Lo, S. (2012). Not all skilled readers have cracked the code: Individual differences in masked form priming. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 38, 152–163. doi:10.1037/a0024953

Bates, D., Mächler, M., Bolker, B., & Walker, S. (2015). Fitting linear mixed-effects models using lme4. *Journal of Statistical Software*, 67, 1–48. doi:10.18637/jss.v067.i01

Bates, D., Maechler, M., Bolker, B., & Walker, S. (2019). *lme4: Linear Mixed-Effects Models using 'Eigen' and S4 (1.1-20)* [software]. Retrieved from <https://CRAN.R-project.org/package=lme4>

Brem, S., Hunkeler, E., Mächler, M., Kronschnabel, J., Karipidis, I. I., Pleisch, G., & Brandeis, D. (2018). Increasing expertise to a novel script modulates the visual N1 ERP in healthy adults. *International Journal of Behavioral Development*, 42, 333–341. doi:10.1177/0165025417727871

Brysbaert, M., & Stevens, M. (2018). Power analysis and effect size in mixed effects models: A Tutorial. *Journal of Cognition*, 1, 9. doi:10.5334/joc.10

- Bürkner, P. C. (2019). Bayesian Item Response Modelling in R with brms and Stan. *arXiv preprint arXiv:1905.09501*.
- Bürkner, P. C. (2021). Bayesian Item Response Modeling in R with brms and Stan. *Journal of Statistical Software*, 100, 1–54. doi:10.18637/jss.v100.i05.
- Bürkner, P.C. (2016). brms: An R package for Bayesian multilevel models using Stan. *Journal of Statistical Software*, 80, 1–28.
- Cassar, M., & Treiman, R. (1997). The beginnings of orthographic knowledge: Children's knowledge of double letters in words. *Journal of Educational Psychology*, 89, 631–644. doi:10.1037/0022-0663.89.4.631
- Castles, A., Davis, C., Cavalot, P., & Forster, K. (2007). Tracking the acquisition of orthographic skills in developing readers: Masked priming effects. *Journal of Experimental Child Psychology*, 97, 165–182. doi: 10.1016/j.jecp.2007.01.006
- Chanceaux, M., & Grainger, J. (2012). Serial position effects in the identification of letters, digits, symbols, and shapes in peripheral vision. *Acta Psychologica*, 141, 149–158. doi:10.1016/j.actpsy.2012.08.001
- Chanceaux, M., Mathôt, S., & Grainger, J. (2014). Effects of number, complexity, and familiarity of flankers on crowded letter identification. *Journal of Vision*, 14, 1–17. doi:10.1167/14.6.7
- Chetail, F. (2017). What do we do with what we learn? Statistical learning of orthographic regularities impacts written word processing. *Cognition*, 163, 103–120. doi:10.1016/j.cognition.2017.02.015
- Chetail, F., & Sauval, K. (2022). Diversity matters: The sensitivity to sublexical orthographic regularities increases with contextual diversity. *Psychonomic Bulletin & Review*, 1–14. doi:10.3758/s13423-021-02029-1
- Dandurand, F., Grainger, J., & Dufau, S. (2010). Learning location-invariant orthographic representations for printed words. *Connection Science*, 22, 25–42. doi: 10.1080/09540090903085768

- Davis, C. J. (2010). The spatial coding model of visual word identification. *Psychological Review*, 117, 713–758. doi:10.1037/a0019738
- Dehaene, S. (2009). *Reading in the brain*. New York: Penguin. doi:10.1111/ijal.12055
- Dehaene, S., & Cohen, L. (2011). The unique role of the visual word form area in reading. *Trends in Cognitive Sciences*, 15, 254–262. doi:10.1016/j.tics.2011.04.003
- Dehaene, S., Cohen, L., Morais, J., & Kolinsky, R. (2015). Illiterate to literate: Behavioural and cerebral changes induced by reading acquisition. *Nature Reviews Neuroscience*, 16, 234–244. doi:10.1038/nrn3924
- Dehaene, S., Cohen, L., Sigman, M., & Vinckier, F. (2005). The neural code for written words: a proposal. *Trends in Cognitive Sciences*, 9, 335–341. doi:10.1016/j.tics.2005.05.004
- Doignon-Camus, N., & Zagar, D. (2014). The syllabic bridge: the first step in learning spelling-to-sound correspondences. *Journal of Child Language*, 41, 1147–1165. doi:10.1017/S0305000913000305
- Duñabeitia, J. A., Dimitropoulou, M., Grainger, J., Hernández, J. A., & Carreiras, M. (2012). Differential sensitivity of letters, numbers, and symbols to character transpositions. *Journal of Cognitive Neuroscience*, 24, 1610–1624. doi:1162/jocn_a_00180
- Duñabeitia, J. A., Lallier, M., Paz-Alonso, P. M., & Carreiras, M. (2015). The impact of literacy on position uncertainty. *Psychological Science*, 26, 548–550. doi:10.1177/0956797615569890
- Duñabeitia, J. A., Orihuela, K., & Carreiras, M. (2014). Orthographic coding in illiterates and literates. *Psychological Science*, 25, 1275–1280. doi:10.1177/0956797614531026
- Estes, W. K. (1975). The locus of inferential and perceptual processes in letter identification. *Journal of Experimental Psychology: General*, 104, 122–145. doi:10.1037/0096-3445.104.2.122

- Faulkenberry, T. J., Ly, A., & Wagenmakers, E. J. (2020). Bayesian inference in numerical cognition: A tutorial using JASP. *Journal of Numerical Cognition*, 6, 231–259. doi:10.5964/jnc.v6i2.288
- Fernández-López, M., Marcet, A., & Perea, M. (2021). Does orthographic processing emerge rapidly after learning a new script? *British Journal of Psychology*, 112, 52–91. doi:10.1111/BJOP.12469
- Forster, K. I., & Davis, C. (1984). Repetition priming and frequency attenuation in lexical access. *Journal of experimental psychology: Learning, Memory, and Cognition*, 10, 680–698. doi:10.1037/0278-7393.10.4.680
- Forster, K. I., & Forster, J. C. (2003). DMDX: A Windows display program with millisecond accuracy. *Behavior Research Methods, Instruments, & Computers*, 35, 116–124. doi:10.3758/BF03195503
- Frankish, C., & Barnes, L. (2008). Lexical and sublexical processes in the perception of transposed-letter anagrams. *Quarterly Journal of Experimental Psychology*, 61, 381–391. doi:10.1080/17470210701664880
- Frankish, C., & Turner, E. (2007). SIHGT and SUNOD: The role of orthography and phonology in the perception of transposed letter anagrams. *Journal of Memory and Language*, 56, 189–211. doi:10.1016/j.jml.2006.11.002
- Friedmann, N., & Gvion, A. (2001). Letter position dyslexia. *Cognitive Neuropsychology*, 18, 673–696. doi:10.1080/02643290143000051
- García-Orza, J., Perea, M., & Muñoz, S. (2010). Are transposition effects specific to letters? *The Quarterly Journal of Experimental Psychology*, 63, 1603–1618. doi:10.1080/17470210903474278
- Gomez, P., & Perea, M. (2020). Masked identity priming reflects an encoding advantage in developing readers. *Journal of Experimental Child Psychology*, 199, 104911. doi:10.1016/j.jecp.2020.104911
- Gómez, P., Marcet, A., & Perea, M. (2021). Are better young readers more likely to confuse their mother with their mother? *Quarterly Journal of*

- Experimental Psychology*, 74, 1542–1552.
doi:10.1177/17470218211012960
- Gomez, P., Ratcliff, R., & Perea, M. (2008). The overlap model: A model of letter position coding. *Psychological Review*, 115, 577–600.
doi:10.1037/a0012667
- Gori, S. y Facoetti, A. (2014). Perceptual learning as a possible new approach for remediation and prevention of developmental dyslexia. *Visual Research*, 99, 78–87. doi:10.1016/j.visres.2013.11.011
- Grainger, J. (2008). Cracking the orthographic code: An introduction. *Language and Cognitive Processes*, 23, 1–35. doi:10.1080/01690960701578013
- Grainger, J. (2018). Orthographic processing: A ‘mid-level’ vision of reading. *Quarterly Journal of Experimental Psychology*, 71, 335–359. doi:10.1080/17470218.2017.1314515
- Grainger, J. (2018). Orthographic processing: A ‘mid-level’ vision of reading: The 44th Sir Frederic Bartlett Lecture. *Quarterly Journal of Experimental Psychology*, 71, 335–359. doi:10.1080/17470218.2017.1314515
- Grainger, J., & Hannagan, T. (2014). What is special about orthographic processing? *Written Language & Literacy*, 17, 225–252. doi:10.1075/wll.17.2.03gra
- Grainger, J., & van Heuven, W. J. B. (2004). Modeling letter position coding in printed word perception. In P. Bonin (Ed.), *Mental lexicon: Some words to talk about words* (pp. 1–23). Hauppauge, NY, US: Nova Science Publishers.
- Grainger, J., & Ziegler, J. (2011). A dual-route approach to orthographic processing. *Frontiers in psychology*, 2, 54. 10.3389/fpsyg.2011.00054
- Grainger, J., Bertrand, D., Lété, B., Beyersmann, E., & Ziegler, J. C. (2016). A developmental investigation of the first-letter advantage. *Journal of Experimental Child Psychology*, 152, 161–172. doi:10.1016/j.jecp.2016.07.016

- Grainger, J., Lété, B., & Bertrand, D. D. S. & Ziegler, J. (2012). Evidence for multiple routes in learning to read. *Cognition*, 123, 280–292. doi:10.1016/j.cognition.2012.01.003
- Grainger, J., Rey, A., & Dufau, S. (2008). Letter perception: From pixels to pandemonium. *Trends in Cognitive Science*, 12, 381–387. doi:10.1016/j.tics.2008.06.006
- Grainger, J., Tydgat, I., & Isselé, J. (2010). Crowding affects letters and symbols differently. *Journal of Experimental Psychology: Human Perception and Performance*, 36, 673–688. doi:10.1037/a0016888
- Guerrera, C., & Forster, K. (2008). Masked form priming with extreme transposition. *Language and Cognitive Processes*, 23, 117–142. doi:10.1080/01690960701579722
- Güven, S., & Friedmann, N. (2019). Developmental Letter Position Dyslexia in Turkish, a Morphologically Rich and Orthographically Transparent Language. *Frontiers in Psychology*, 10. doi:10.3389/fpsyg.2019.02401
- Jackson, N., & Coltheart, M. (2001). Routes to reading success and failure. Hove, UK: Psychology Press.
- Jeffreys, H. (1961). Theory of probability (3rd ed.). Oxford: Oxford University Press, Clarendon Press.
- Kohnen, S., Nickels, L., Castles, A., Friedmann, N., and McArthur, G. (2012). When 'slime' becomes 'smile': developmental letter position dyslexia in English. *Neuropsychologia*, 50, 3681–3692. doi:10.1016/j.neuropsychologia.2012.07.016
- Krueger, L. E. (1978). A theory of perceptual matching. *Psychological Review*, 85, 278–304. doi:10.1037//0033-295X.85.4.278
- Ktori, M., Bertrand, D., & Grainger, J. (2019). What's special about orthographic processing? Further evidence from transposition effects in same-different matching. *Quarterly Journal of Experimental Psychology*, 72, 1780–1789. doi:10.1177/1747021818811448

- Kuznetsova, A., Brockhoff, P. B., & Christensen, R. H. B. (2017). lmerTest package: Tests in linear mixed effects models. *Journal of Statistical Software*, 82, 1–26. doi:10.18637/jss.v082.i13
- Lachmann, T., & van Leeuwen, C. (2014). Reading as functional coordination: not recycling but a novel synthesis. *Frontiers in psychology*, 5, 1046. doi:10.3389/fpsyg.2014.01046
- Lally, C., Taylor, J. S. H., Lee, C. H., & Rastle, K. (2020). Shaping the precision of letter position coding by varying properties of a writing system. *Language, Cognition and Neuroscience*, 35, 374–382. doi:10.1080/23273798.2019.1663222
- Leininger, M., Myslín, M., Rayner, K., & Levy, R. (2017). Do resource constraints affect lexical processing? Evidence from eye movements. *Journal of Memory and Language*, 93, 82–103. doi:10.1016/j.jml.2016.09.002
- Lelonkiewicz, J. R., Ktori, M., & Crepaldi, D. (2020). Morphemes as letter chunks: Discovering affixes through visual regularities. *Journal of Memory and Language*, 115, 104152. doi:10.1016/j.jml.2020.104152
- Logan, G. D. (2021). Serial order in perception, memory, and action. *Psychological Review*, 128, 1–44. doi:10.1037/rev0000253
- Lupker, S. J., Perea, M., & Davis, C. J. (2008). Transposed-letter effects: Consonants, vowels and letter frequency. *Language and Cognitive Processes*, 23, 93–116. doi:10.1080/01690960701579714
- Mano, Q. R., & Kloos, H. (2018). Sensitivity to the regularity of letter patterns within print among preschoolers: Implications for emerging literacy. *Journal of Research in Childhood Education*, 32, 379–391. doi:10.1080/02568543.2018.1497736
- Marcet, A., Perea, M., Baciero, A., & Gomez, P. (2019). Can letter position encoding be modified by visual perceptual elements? *Quarterly Journal of Experimental Psychology*, 72, 1344–1353. doi:10.1177/1747021818789876

- Marinus, E., Kezilas, Y., Kohnen, S., Robidoux, S., & Castles, A. (2018). Who are the noisiest neighbors in the hood? Using error analyses to study the acquisition of letter-position processing. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, *44*, 1384–1396. doi:10.1037/xlm0000524
- Massol, S., Duñabeitia, J. A., Carreiras, M., & Grainger, J. (2013). Evidence for Letter-Specific Position Coding Mechanisms. *PLOS ONE*, *8*, e68460. doi:10.1371/journal.pone.0068460
- Maurer, U., & McCandliss, B. D. (2007). The development of visual expertise for words: the contribution of electrophysiology. In Grigorenko, E. L., & Naples, A.J., (Eds.) *Single-Word Reading: Cognitive, Behavioral and Biological Perspectives* (pp. 43–64). Mahwah, NJ: Lawrence Erlbaum Associates.
- Maurer, U., Blau, V. C., Yoncheva, Y. N., & McCandliss, B. D. (2010). Development of visual expertise for reading: rapid emergence of visual familiarity for an artificial script. *Developmental Neuropsychology*, *35*, 404–422. doi:10.1080/87565641.2010.480916
- Maurer, U., Brandeis, D., & McCandliss, B. D. (2005). Fast, visual specialization for reading in English revealed by the topography of the N170 ERP response. *Behavioral and Brain Functions*, *1*, 13. doi:10.1186/1744-9081-1-13
- Maurer, U., Brem, S., Bucher, K., & Brandeis, D. (2005). Emerging neurophysiological specialization for letter strings. *Journal of Cognitive Neuroscience*, *17*, 1532–1552. doi:10.1162/089892905774597218
- Maurer, U., Brem, S., Kranz, F., Bucher, K., Benz, R., Halder, P., ... & Brandeis, D. (2006). Coarse neural tuning for print peaks when children learn to read. *Neuroimage*, *33*, 749–758. doi:10.1016/j.neuroimage.2006.06.025
- McCandliss, B. D., & Noble, K. G. (2003). The development of reading impairment: a cognitive neuroscience model. *Mental retardation and*

- developmental disabilities research reviews*, 9, 196–205.
doi:10.1002/mrdd.10080
- McCann, R. S., Folk, C. L., & Johnston, J. C. (1992). The role of spatial attention in visual word processing. *Journal of Experimental Psychology: Human Perception and Performance*, 18, 1015–1029. doi:10.1037/0096-1523.18.4.1015
- Morey, R.D., & Rouder, J.N., (2018). Package 'BayesFactor'. *R Package version 0.9.12-4.2*. Available on: <https://richarddmorey.github.io/BayesFactor/>
- Muñoz, S., Perea, M., García-Orza, J., & Barber, H. A. (2012). Electrophysiological signatures of masked transposition priming in a same-different task: Evidence with strings of letters vs. pseudoletters. *Neuroscience Letters*, 515, 71–76. doi:10.1016/j.neulet.2012.03.021
- Norris, D., & Kinoshita, S. (2012). Reading through a noisy channel: Why there's nothing special about the perception of orthography. *Psychological Review*, 119, 517–545. doi:10.1037/a0028450
- Norris, D., Kinoshita, S., & van Casteren, M. (2010). A stimulus sampling theory of letter identity and order. *Journal of Memory and Language*, 62, 254–271. doi:10.1016/j.jml.2009.11.002
- O'Brien, B. A. (2014). The development of sensitivity to sublexical orthographic constraints: An investigation of positional frequency and consistency using a wordlikeness choice task. *Reading Psychology*, 35, 285–311. doi:10.1080/02702711.2012.724042
- Pacton, S., Perruchet, P., Fayol, M., & Cleeremans, A. (2001). Implicit learning out of the lab: The case of orthographic regularities. *Journal of Experimental Psychology: General*, 130, 401–426. doi:10.1037/0096-3445.130.3.401
- Pagán, A., Blythe, H. I., & Liversedge, S. P. (2021). The influence of children's reading ability on initial letter position encoding during a reading-like task.

- Journal of Experimental Psychology: Learning, Memory, and Cognition*, 47, 1186–1203. doi:10.1037/xlm00 00989
- Perea, M., & Carreiras, M. (2008). Do orthotactics and phonology constrain the transposed-letter effect? *Language and Cognitive Processes*, 23, 69–92. doi:10.1080/01690960701578146
- Perea, M., & Lupker, S. J. (2003). Does jugde activate COURT? Transposed-letter similarity effects in masked associative priming. *Memory and Cognition*, 31, 829–841. doi:10.3758/BF03196438
- Perea, M., & Lupker, S. J. (2004). Can CANISO activate CASINO? Transposed-letter similarity effects with nonadjacent letter positions. *Journal of Memory and Language*, 51, 231–246. doi:10.1016/j.jml.2004.05.005
- Perea, M., Jiménez, M., & Gomez, P. (2016). Does location uncertainty in letter position coding emerge because of literacy training? *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 42, 996–1001. doi:10.1037/xlm0000208
- Perea, M., Jiménez, M., Suárez-Coalla, P., Fernández, N., Viña, C. y Cuetos, F. (2014). Ability for voice recognition is a marker for dyslexia in children. *Experimental Psychology*, 61, 480–487. doi:10.1027/1618-3169/a000265
- Perea, M., Marcet, A., & Gomez, P. (2016). How do Scrabble players encode letter position during reading? *Psicothema*, 28, 7–12. doi:10.7334/psicothema2015.167
- Perfetti, C. A. (2017). Lexical quality revisited. In E. Segers & P. van den Broek (Eds.), *Developmental perspectives in written language and literacy: In honor of Ludo Verhoeven* (pp. 51–68). Amsterdam, Netherlands: John Benjamins.
- Pleisch, G., Karipidis, I. I., Brauchli, C., Röthlisberger, M., Hofstetter, C., Stämpfli, P., ... & Brem, S. (2019). Emerging neural specialization of the ventral occipitotemporal cortex to characters through phonological association

- learning in preschool children. *NeuroImage*, 189, 813–831. doi: 10.1016/j.neuroimage.2019.01.046
- Price, C. J., & Devlin, J. T. (2011). The interactive account of ventral occipitotemporal contributions to reading. *Trends in Cognitive Sciences*, 15, 246–253. doi: 10.1016/j.tics.2011.04.001
- R Core Team (2019). *R: A language and environment for statistical computing*. Vienna: R Foundation for Statistical Computing. Retrieved from <http://www.R-project.org/>
- R Core Team (2021). *R: A language and environment for statistical computing*. Viena: R Foundation for Statistical Computing. Retrieved from <https://www.R-project.org/>
- Ratcliff, R. (1981). A theory of order relations in perceptual matching. *Psychological Review*, 88, 552–572. doi:10.1037/0033-295X.88.6.552
- Reicher, G. M. (1969). Perceptual recognition as a function of meaningfulness of stimulus material. *Journal of Experimental Psychology*, 81, 275–280. doi:10.1037/h0027768
- Rouder, J. N., Speckman, P. L., Sun, D., Morey, R. D., & Iverson, G. (2009). Bayesian t tests for accepting and rejecting the null hypothesis. *Psychonomic Bulletin & Review*, 16, 225–237. doi:10.3758/PBR.16.2.225
- Rumelhart, D. E. (1985). Toward an interactive model of reading. In S. Dornic (Ed.), *Attention and performance VI* (pp. 573–603). Erlbaum.
- Saygin, Z. M., Osher, D. E., Norton, E. S., Youssoufian, D. A., Beach, S. D., Feather, J., ... & Kanwisher, N. (2016). Connectivity precedes function in the development of the visual word form area. *Natural Neuroscience*, 19, 1250–1255. doi:10.1038/nn.4354
- Scaltritti, M., & Balota, D. A. (2013). Are all letters really processed equally and in parallel? Further evidence of a robust first letter advantage. *Acta Psychologica*, 144, 397–410. doi:10.1016/j.actpsy.2013.07.018

- Scaltritti, M., Dufau, S., & Grainger, J. (2018). Stimulus orientation and the first-letter advantage. *Acta Psychologica*, 183, 37–42. doi:10.1016/j.actpsy.2017.12.009
- Schönbrodt, F. D., & Wagenmakers, E. J. (2018). Bayes factor design analysis: Planning for compelling evidence. *Psychonomic Bulletin & Review*, 25, 128–142. doi:10.3758/s13423-017-1230-y
- Schubert, T., Badcock, N., & Kohnen, S. (2017). Development of children's identity and position processing for letter, digit, and symbol strings: A cross-sectional study of the primary school years. *Journal of Experimental Child Psychology*, 162, 163–180. doi:10.1016/j.jecp.2017.05.008
- Sellés, P., Martínez, T., Vidal-Abarca, E., & Gilabert, R. (2008). *BIL 3–6. Bateria de Inicio a la Lectura* [Battery to Assess the Abilities Related with Early Reading Acquisition]. Madrid, Spain: Instituto de Ciencias de la Educación.
- Shetreet, E., & Friedmann, N. (2011). Induced letter migrations between words and what they reveal about the orthographic-visual analyzer. *Neuropsychologia*, 49, 339–351. doi:10.1016/j.neuropsychologia.2010.11.026
- Singmann, H., & Gronau, Q. F. (2021). *Bayes Factors for brms models* [Computer software]. doi:10.5281/zenodo.4904827
- Staub, A., & Goddard, K. (2019). The role of preview validity in predictability and frequency effects on eye movements in reading. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 45, 110–127. doi:10.1037/xlm0000561
- Taylor, J. S. H., Davis, M. H., & Rastle, K. (2017). Comparing and validating methods of reading instruction using behavioural and neural findings in an artificial orthography. *Journal of Experimental Psychology: General*, 146, 826–858. doi:10.1037/xge0000301

- Taylor, J. S. H., Plunkett, K., & Nation, K. (2011). The influence of consistency, frequency, and semantics on learning to read: An artificial orthography paradigm. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 37, 60–76. doi:10.1037/a0020126
- Tydgat, I., & Grainger, J. (2009). Serial position effects in the identification of letters, digits, and symbols. *Journal of Experimental Psychology: Human Perception and Performance*, 35, 480–498. doi:10.1037/a0013027
- Vejnović, D., & Zdravković, S. (2015). Side flankers produce less crowding, but only for letters. *Cognition*, 143, 217–227. doi:10.1016/j.cognition.2015.07.003
- Vidal, C., Content, A., & Chetail, F. (2017). BACS: The Brussels Artificial Character Sets for studies in cognitive psychology and neuroscience. *Behavior Research Methods*, 49, 2093–2112. doi:10.3758/s13428-016-0844-8
- Vidal, Y., Viviani, E., Zoccolan, D., & Crepaldi, D. (2021). A general-purpose mechanism of visual feature association in visual word identification and beyond. *Current Biology*, 31, 1261–1267. doi:10.1016/j.cub.2020.12.017
- Wagenmakers, E.-J., Marsman, M., Jamil, T., Ly, A., Verhagen, J., Love, J., ... Morey, R. D. (2017). Bayesian inference for psychology. Part I: Theoretical advantages and practical ramifications. *Psychonomic Bulletin & Review*, 25, 35–57. doi:10.3758/s13423-017-1343-3
- Wagenmakers, E.J., Love, J., Marsman, M., Jamil, T., Ly, A., Verhagen, J., ...& Meerhoff, F. (2018). Bayesian inference for psychology. Part II: Example applications with JASP. *Psychonomic Bulletin & Review*, 25, 58–76. doi:10.3758/s13423-017-1323-7
- Wheeler, D. D. (1970). Processes in word recognition. *Cognitive Psychology*, 1, 59–85. doi:10.1016/0010-0285(70)90005-8
- Whitney, C. (2001). How the brain encodes the order of letters in a printed word: The SERIOL model and selective literature review. *Psychonomic Bulletin & Review*, 8, 221–243. doi:10.3758/BF03196158

- Whitney, C., Bertrand, D., & Grainger, J. (2011). On coding the position of letters in words: a test of two models. *Experimental Psychology*, 59, 109–114. doi:10.1027/1618-3169/a000132
- Winkel, H., Perea, M., & Peart, E. (2014). Testing the flexibility of the modified receptive field (MRF) theory: Evidence from an unspaced orthography (Thai). *Acta Psychologica*, 150, 55–60. doi:10.1016/j.actpsy.2014.04.008
- Ziegler, J. C., Bertrand, D., L  t  , B., & Grainger, J. (2014). Orthographic and phonological contributions to reading development: Tracking developmental trajectories using masked priming. *Developmental Psychology*, 50, 1026-1036. doi:10.1037/a0035187
- Ziegler, J. C., Pech-Georgel, C., Dufau, S., & Grainger, J. (2010). Rapid processing of letters, digits and symbols: what purely visual-attentional deficit in developmental dyslexia? *Developmental Science*, 13, F8–F14. doi:10.1111/j.1467-7687.2010.00983.x

The resilience of letter detectors

- Angele, B., Slattery, T. J., & Rayner, K. (2016). Two stages of parafoveal processing during reading: Evidence from a display change detection task. *Psychonomic Bulletin & Review*, 23, 1241–1249. doi:10.3758/s13423-015-0995-0
- Angele, B., Tran, R., & Rayner, K. (2013). Parafoveal–foveal overlap can facilitate ongoing word identification during reading: Evidence from eye movements. *Journal of Experimental Psychology: Human Perception and Performance*, 39, 526–538. doi:10.1037/a0029492
- Aschenbrenner, A. J., & Yap, M. J. (2019). The influence of relatedness proportion on the joint relationship among word frequency, stimulus quality, and semantic priming in the lexical decision task. *Quarterly Journal of Experimental Psychology*, 72, 2452–2461. doi:10.1177/1747021819845317

- Barnhart, A. S., & Goldinger, S. D. (2013). Rotation reveals the importance of configural cues in handwritten word perception. *Psychonomic Bulletin & Review*, 20, 1319–1326. doi:10.3758/s13423-013-0435-y
- Bates, D., Mächler, M., Bolker, B., & Walker, S. (2015). Fitting linear mixed-effects models using lme4. *Journal of Statistical Software*, 67, 1–48. doi:10.18637/jss.v067.i01
- Bates, D., Mächler, M., Bolker, B., & Walker, S. (2019). *lme4: Linear Mixed-Effects Models using 'Eigen' and S4 (1.1-20)* [software]. Retrieved from <https://CRAN.R-project.org/package=lme4>
- Bentin, S., Mouchetant-Rostaing, Y., Giard, M. H., Echallier, J. F., & Pernier, J. (1999). ERP manifestations of processing printed words at different psycholinguistic levels: time course and scalp distribution. *Journal of Cognitive Neuroscience*, 11, 235–260. doi:10.1162/089892999563373
- Benyhe, A., & Csibri, P. (2021). Can rotated words be processed automatically? Evidence from rotated repetition priming. *Memory & Cognition*, 49, 1163–1171. doi:10.3758/s13421-021-01147-4
- Blythe, H. I., Juhasz, B. J., Tbaily, L. W., Rayner, K., & Liversedge, S. P. (2019). Reading sentences of words with rotated letters: An eye movement study. *Quarterly Journal of Experimental Psychology*, 72(7), 1790–1804. doi:10.1177/1747021818810381
- Brysbaert, M., & New, B. (2009). Moving beyond Kučera and Francis: A critical evaluation of current word frequency norms and the introduction of a new and improved word frequency measure for American English. *Behavior Research Methods*, 41, 977–990. doi:10.3758/BRM.41.4.977
- Bürkner, P. (2021). Bayesian Item Response Modeling in R with brms and Stan. *Journal of Statistical Software*, 100, 1–54. doi:10.18637/jss.v100.i05.
- Carreiras, M., Armstrong, B. C., Perea, M., & Frost, R. (2014). The what, when, where, and how of visual word recognition. *Trends in Cognitive Sciences*, 18, 90–98.

- Carreiras, M., Duñabeitia, J. A., & Molinaro, N. (2009). Consonants and vowels contribute differently to visual word recognition: ERPs of relative position priming. *Cerebral Cortex*, *19*, 2659–2670. doi:10.1093/cercor/bhp019
- Carreiras, M., Vergara, M., Perea, M. (2007). ERP correlates of transposed-letter similarity effects: are consonants processed differently from vowels? *Neuroscience Letters*, *419*, 219–224. doi:10.1016/j.neulet.2007.04.053.
- Chauncey, K., Holcomb, P. J., & Grainger, J. (2008). Effects of stimulus font and size on masked repetition priming: An event-related potentials (ERP) investigation. *Language and Cognitive Processes*, *23*, 183–200. doi:10.1080/01690960701579839
- Clifton Jr, C., Ferreira, F., Henderson, J. M., Inhoff, A. W., Liversedge, S. P., Reichle, E. D., & Schotter, E. R. (2016). Eye movements in reading and information processing: Keith Rayner's 40 year legacy. *Journal of Memory and Language*, *86*, 1–19. doi:10.1016/j.jml.2015.07.004
- Cohen, L., & Dehaene, S. (2009). Ventral and dorsal contributions to word reading. In M. S. Gazzaniga (Eds), *The cognitive neuroscience* (pp. 789–804). Cambridge, MA: MIT Press.
- Cohen, L., Dehaene, S., Vinckier, F., Jobert, A., & Montavont, A. (2008). Reading normal and degraded words: Contribution of the dorsal and ventral visual pathways. *Neuroimage*, *40*, 353–366. doi:10.1016/j.neuroimage.2007.11.036
- Dehaene, S., Cohen, L., Sigman, M., & Vinckier, F. (2005). The neural code for written words: A proposal. *Trends in Cognitive Sciences*, *9*, 335–341. doi:10.1016/j.tics.2005.05.004
- Dickson, D. S., & Federmeier, K. D. (2014). Hemispheric differences in orthographic and semantic processing as revealed by event-related potentials. *Neuropsychologia*, *64*, 230–239. doi:10.1016/j.neuropsychologia.2014.09.037

- Duchon, A., Perea, M., Sebastián-Gallés, N., Martí, A., & Carreiras, M. (2013). EsPal: One-stop shopping for Spanish word properties. *Behavior Research Methods*, 45, 1246–1258. doi:10.3758/s13428-013-0326-1
- Dufau, S., Grainger, J., & Holcomb, P. J. (2008). An ERP investigation of location invariance in masked repetition priming. *Cognitive, Affective, & Behavioral Neuroscience*, 8, 222–228. doi:10.3758/CABN.8.2.222
- Engbert, R., Nuthmann, A., Richter, E. M., & Kliegl, R. (2005). SWIFT: A Dynamical Model of Saccade Generation During Reading. *Psychological Review*, 112(4), 777–813. doi:10.1037/0033-295X.112.4.777
- Fernández-López, M., Marcet, A., & Perea, M. (2019). Can response congruency effects be obtained in masked priming lexical decision? *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 45, 1683–1702. doi:10.1037/xlm0000666
- Fernández-López, M., Mirault, J., Grainger, J., & Perea, M. (2021). How resilient is reading to letter rotations? A parafoveal preview investigation. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 47, 2029–2042. doi: 10.1037/xlm0000999
- Forster, K. I., & Davis, C. (1984). Repetition priming and frequency attenuation in lexical access. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 10, 680–698. doi:10.1037/0278-7393.10.4.680
- Forster, K. I., & Forster, J. C. (2003). DMDX: A Windows display program with millisecond accuracy. *Behavior Research Methods, Instruments, & Computers*, 35, 116–124. doi:10.3758/BF03195503
- Fox, J., Weisberg, S. (2019). *An R Companion to Applied Regression*. R package version 3.0-8. Retrieved from <https://CRAN.R-project.org/package=car>
- Gomez, P., & Perea, M. (2014). Decomposing encoding and decisional components in visual-word recognition: A diffusion model analysis. *Quarterly Journal of Experimental Psychology*, 67(12), 2455–2466. doi:10.1080/17470218.2014.937447

- Grainger, J. (2018). Orthographic processing: A 'mid-level' vision of reading. *Quarterly Journal of Experimental Psychology*, 71, 335–359. doi:10.1080/17470218.2017.1314515
- Grainger, J., & Dufau, S. (2012). The front-end of visual word recognition, in J.S. Adelman (Ed). *Visual word recognition, Vol. 1* (pp. 159–184). Hove: Psychology Press.
- Grainger, J., & Holcomb, P. J. (2009). Watching the word go by: On the time-course of component processes in visual word recognition. *Language and Linguistics Compass*, 3, 128–156. doi:10.1111/j.1749-818X.2008.00121.x
- Greenhouse, S. W., & Geisser, S. (1959). On methods in the analysis of profile data. *Psychometrika*, 24, 95–112. doi:10.1007/BF02289823
- Gutiérrez-Sigut, E., Vergara-Martínez, M., & Perea, M. (2019). Deaf readers benefit from lexical feedback during orthographic processing. *Scientific Reports*, 9, 12321. doi:10.1038/s41598-019-48702-3
- Hannagan, T., Ktori, M., Chanceaux, M., & Grainger, J. (2012). Deciphering CAPTCHAs: What a Turing Test Reveals about Human Cognition. *PLOS ONE*, 7, e32121. doi:10.1371/journal.pone.0032121
- Holcomb, P. J., & Grainger, J. (2006). On the time course of visual word recognition: An event-related potential investigation using masked repetition priming. *Journal of Cognitive Neuroscience*, 18, 1631–1643. doi:10.1162/jocn.2006.18.10.1631
- Holcomb, P. J., & Grainger, J. (2007). Exploring the temporal dynamics of visual word recognition in the masked repetition priming paradigm using event-related potentials. *Brain Research*, 1180, 39–58. doi:10.1016/j.neuropsychologia.2021.107924
- Jung, T. P., Makeig, S., Westerfield, M., Townsend, J., Courchesne, E., & Sejnowski, T. J. (2000). Removal of eye activity artifacts from visual event-related potentials in normal and clinical subjects. *Clinical Neurophysiology*, 111, 1745–1758. doi:10.1016/S1388-2457(00)00386-2

- Keuleers, E., & Brysbaert, M. (2010). Wuggy: A multilingual pseudoword generator. *Behavior Research Methods*, 42, 627–633. doi:10.3758/BRM.42.3.627
- Kim, A. E., & Straková, J. (2012). Concurrent effects of lexical status and letter-rotation during early stage visual word recognition: Evidence from ERPs. *Brain Research*, 1468, 52–62. doi:10.1016/j.brainres.2012.04.008
- Kiyonaga, K., Grainger, J., Midgley, K., & Holcomb, P. J. (2007). Masked cross-modal repetition priming: An event-related potential investigation. *Language and Cognitive Processes*, 22, 337–376. doi: 10.1080/01690960600652471
- Koriat, A., & Norman, J. (1984). What is rotated in mental rotation? *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 10, 421–434. doi:10.1037/0278-7393.10.3.421
- Kuznetsova, A., Brockhoff, P. B., & Christensen, R. H. B. (2016). *lmerTest: Tests for random and fixed effects for linear mixed effect models (lmer objects of lme4 package)*. R package Version 2.0 –33. Retrieved from <http://CRAN.Rproject.org/package=lmerTest>
- Lenth, R., Singmann, H., Love, J., Buerkner, P., & Herve, M. (2020). *emmeans: Estimated Marginal Means*. R package version 1.4.8. Retrieved from <https://CRAN.R-project.org/package=emmeans>
- Logothetis, N. K., & Pauls, J. (1995). Psychophysical and physiological evidence for viewer-centered object representations in the primate. *Cerebral Cortex*, 5, 270–288. doi:10.1093/cercor/5.3.270
- Mathôt, S., Schreij, D., & Theeuwes, J. (2012). OpenSesame: An open-source, graphical experiment builder for the social sciences. *Behavior Research Methods*, 44, 314–324. doi:10.3758/s13428-011-0168-7
- Maurer, U., Brem, S., Bucher, K., & Brandeis, D. (2005). Emerging neurophysiological specialization for letter strings. *Journal of Cognitive Neuroscience*, 17, 1532–1552. doi:10.1162/089892905774597218
- Mayall, K., & Humphreys, G. W. (1996). Case mixing and the task-sensitive disruption of lexical processing. *Journal of Experimental Psychology*:

Learning, Memory, and Cognition, 22, 278–294. doi:10.1037/0278-7393.22.2.278

Morey, R.D., & Rouder, J.N., (2018). Package ‘BayesFactor’. *R Package version 0.9.12-4.2*. Retrieved from <https://richarddmores.github.io/BayesFactor/>

New, B., Pallier, C., Brysbaert, M., & Ferrand, L. (2004). Lexique 2: A new French lexical database. *Behavior Research Methods, Instruments, & Computers*, 36, 516–524. doi:10.3758/BF03195598

Oldfield, R. C. (1971). The assessment and analysis of handedness: the Edinburgh inventory. *Neuropsychologia*, 9, 97–113. doi:10.1016/0028-3932(71)90067-4

Peirce, J. W., & Macaskill, M. R. (2018). *Building experiments in PsychoPy*. Sage.

Perea, M., Fernández-López, M., & Marcet, A. (2020). Does CaSe-MiXinG disrupt the access to lexico-semantic information? *Psychological Research*. 84, 981–989. doi: 10.1007/s00426-018-1111-7

Perea, M., Marcet, A., & Fernández-López, M. (2018). Does letter rotation slow down orthographic processing in word recognition? *Psychonomic Bulletin and Review*, 25, 2295–2300. doi:10.3758/s13423-017-1428-z

Perea, M., Rosa, E., & Marcet, A. (2017). Where is the locus of the lowercase advantage during sentence reading? *Acta Psychologica*, 177, 30–35. doi: 10.1016/j.actpsy.2017.04.007

Perea, M., Vergara-Martínez, M., & Gomez, P. (2015). Resolving the locus of cAsE aLtErNaTiOn effects in visual word recognition: Evidence from masked priming. *Cognition*, 142, 39–43. doi:10.1016/j.cognition.2015.05.007

Perea, M., Vergara-Martínez, M., Marcet, A., Abu Mallouh, R., & Fernández-López, M. (2020). When does rotation disrupt letter encoding? Testing the resilience of letter detectors in the initial moments of processing. *Memory & Cognition*, 48, 704–709. doi:10.3758/s13421-020-01013-9

- Petit, J. P., Midgley, K. J., Holcomb, P. J., & Grainger, J. (2006). On the time course of letter perception: A masked priming ERP investigation. *Psychonomic Bulletin & Review*, 13, 674–681. doi:10.3758/bf03193980
- Pickering, E. C., & Schweinberger, S. R. (2003). N200, N250r, and N400 Event-Related Brain Potentials Reveal Three Loci of Repetition Priming for Familiar Names. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 29, 1298–1311. doi:10.1037/0278-7393.29.6.1298
- Pollatsek, A., & Well, A. D. (1995). On the use of counterbalanced designs in cognitive research: A suggestion for a better and more powerful analysis. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 21, 785–794. doi:10.1037/0278-7393.21.3.785
- Qiao, E., Vinckier, F., Szwed, M., Naccache, L., Valabrègue, R., Dehaene, S., Cohen, L. (2010). Unconsciously deciphering handwriting: subliminal invariance for handwritten words in the visual word form area. *Neuroimage*, 49, 1786–1799. doi:10.1016/j.neuroimage.2009.09.034
- R Core Team (2020). *R: A language and environment for statistical computing*. Vienna: R Foundation for Statistical Computing. Retrieved from <http://www.R-project.org/>
- R Core Team (2021). *R: A language and environment for statistical computing*. R Foundation for Statistical Computing. Retrieved from <https://www.R-project.org/>
- Rayner, K. (1975). Parafoveal identification during a fixation in reading. *Acta Psychologica*, 39, 271–281. doi:10.1016/0001-6918(75)90011-6
- Rayner, K. (2009). The 35th Sir Frederick Bartlett Lecture: Eye movements and attention in reading, scene perception, and visual search. *Quarterly Journal of Experimental Psychology*, 62, 1457–1506. doi:10.1080/17470210902816461
- Rayner, K., Inhoff, A. W., Morrison, R. E., Slowiaczek, M. L., & Bertera, J. H. (1981). Masking of foveal and parafoveal vision during eye fixations in

- reading. *Journal of Experimental Psychology: Human Perception and Performance*, 7, 167–179. doi:10.1037/0096-1523.7.1.167
- Rayner, K., Pollatsek, A., Ashby, J., & Clifton Jr, C. (2012). *Psychology of reading*. New York: Psychology Press. doi:10.4324/9780203155158
- Rayner, K., Well, A. D., Pollatsek, A., & Bertera, J. H. (1982). The availability of useful information to the right of fixation in reading. *Perception & Psychophysics*, 31(6), 537–550. doi:10.3758/BF03204186
- Reichle, E. D., Pollatsek, A., Fisher, D. L., & Rayner, K. (1998). Toward a model of eye movement control in reading. *Psychological Review*, 105, 125–157. doi:10.1037/0033-295x.105.1.125
- Reichle, E. D., Rayner, K., & Pollatsek, A. (2003). The EZ Reader model of eye-movement control in reading: Comparisons to other models. *Behavioral and Brain Sciences*, 26, 445–526. doi:10.1017/S0140525X03290105
- Rossion, B., Joyce, C. A., Cottrell, G. W., & Tarr, M. J. (2003). Early lateralization and orientation tuning for face, word, and object processing in the visual cortex. *Neuroimage*, 20, 1609–1624. doi:10.1016/j.neuroimage.2003.07.010
- Rouder, J. N., Speckman, P. L., Sun, D., Morey, R. D., & Iverson, G. (2009). Bayesian *t* tests for accepting and rejecting the null hypothesis. *Psychonomic Bulletin & Review*, 16(2), 225–237. doi:10.3758/PBR.16.2.225
- Slattery, T.J. (2016). Eye movements: From psycholinguistics to font design. In M. Dyson, (Ed). *Digital fonts and reading*. (pp. 54–78) Singapore: World Scientific. doi:10. 1142/9789814759540_0004
- Snell, J., van Leipsig, S., Grainger, J., & Meeter, M. (2018). OB1-reader: A model of word recognition and eye movements in text reading. *Psychological Review*, 125, 969–984. doi:10.1037/rev0000119
- Tinker, M. A. (1958). Recent studies of eye movements in reading. *Psychological Bulletin*, 55, 215–231. doi:10.1037/h0041228

- Tinker, M. A., & Paterson, D. G. (1931). Studies of typographical factors influencing speed of reading. VII. Variations in color of print and background. *Journal of Applied Psychology*, 15, 471–479. doi:10.1037/h0076001
- van Heuven, W. J., Mandera, P., Keuleers, E., & Brysbaert, M. (2014). SUBTLEX-UK: A new and improved word frequency database for British English. *Quarterly journal of experimental psychology*, 67, 1176–1190. doi:10.1080/17470218.2013.850521
- Vasilev, M. R., Slattery, T. J., Kirkby, J. A., & Angele, B. (2018). What are the costs of degraded parafoveal previews during silent reading? *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 44(3), 371–386. doi: 10.1037/xlm0000433
- Vergara-Martínez, M., Gomez, P., Jiménez, M., & Perea, M. (2015). Lexical enhancement during prime-target integration: ERP evidence from matched-case identity priming. *Cognitive, Affective, and Behavioral Neuroscience*, 15, 492–504. doi:10.3758/s13415-014-0330-7
- Vergara-Martínez, M., Gutierrez-Sigut, E., Perea, M., Gil-López, C., & Carreiras, M. (2021). The time course of processing handwritten words: An ERP investigation. *Neuropsychologia*, 159, 107924. doi:10.1016/j.neuropsychologia.2021.107924
- Vergara-Martínez, M., Perea, M., & Leone-Fernandez, B. (2020). The time course of the lowercase advantage in visual word recognition: An ERP investigation. *Neuropsychologia*, 146, 107556. doi:10.1016/j.neuropsychologia.2020.107556
- Vinckier, F., Naccache, L., Papeix, C., Forget, J., Hahn-Barma, V., Dehaene, S., & Cohen, L. (2006). “What” and “where” in Word Reading: Ventral Coding of Written Words Revealed by Parietal Atrophy. *Journal of Cognitive Neuroscience*, 18, 1998–2012. doi:10.1162/jocn.2006.18.12.1998
- Wagenmakers, E.-J., Marsman, M., Jamil, T., Ly, A., Verhagen, J., Love, J., ... Morey, R. D. (2017). Bayesian inference for psychology. Part I:

- Theoretical advantages and practical ramifications. *Psychonomic Bulletin & Review*, 25(1), 35–57. doi:10.3758/s13423-017-1343-3
- Whitney, C. (2001). How the brain encodes the order of letters in a printed word: The SERIOL model and selective literature review. *Psychonomic Bulletin & Review*, 8(2), 221–243. doi:10.3758/BF03196158
- Whitney, C. (2002). An explanation of the length effect for rotated words. *Cognitive Systems Research*, 3, 113–119. doi:10.1016/S1389-0417(01)00050-X
- Whitney, C. (2010). *One mechanism or two: A commentary on reading normal and degraded words*. Unpublished manuscript. Retrieved from <http://www.cs.umd.edu/~shankar/cwhitney/selected-publications.html>
- Yang, H., & Lupker, S. J. (2019). Does letter rotation decrease transposed letter priming effects? *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 45, 2309–2318 doi:10.1037/xlm0000697
- Yarkoni, T., Balota, D., & Yap, M. (2008). Moving beyond Coltheart's N: A new measure of orthographic similarity. *Psychonomic Bulletin & Review*, 15, 971-979. doi:10.3758/PBR.15.5.971

Methodological contribution to the visual word recognition research

- Adelman, J. S., Johnson, R. L., McCormick, S. F., McKague, M., Kinoshita, S., Bowers, J. S., ... Scaltritti, M. (2014). A behavioral database for masked form priming. *Behavior Research Methods*, 46, 1052–1067. doi:10.3758/s13428-013-0442-y
- Balota, D. A., Yap, M. J., Cortese, M. J., & Watson, J. M. (2008). Beyond mean response latency: Response time distributional analyses of semantic priming. *Journal of Memory and Language*, 59, 495–523. doi:10.1016/j.jml.2007.10.004

- Barr, D. J., Levy, R., Scheepers, C., & Tily, H. J. (2013). Random effects structure for confirmatory hypothesis testing: Keep it maximal. *Journal of Memory and Language*, 68, 255–278. doi:10.1016/j.jml.2012.11.001
- Bates, D., Maechler, M., Bolker, B., & Walker, S. (2015). Fitting linear mixed-effects models using lme4. *Journal of Statistical Software*, 67, 1–48. doi:10.18637/jss.v067.i01
- Bayliss, L., Davis, C. J., Brysbaert, M., Luyten, S., & Rastle, K. (2009). *How are lexical decisions to word targets influenced by unrelated masked primes?* Unpublished manuscript. (Retrieved on December 21, 2017 at <http://crr.ugent.be/papers/Bayliss%20et%20al.%202008.pdf>)
- Bonnet, C., & Dresp, B. (1993). A fast procedure for studying conditional accuracy functions. *Behavior Research Methods, Instruments, & Computers* 25, 2–8 doi:10.3758/BF03204443
- Bürkner, P. C. (2020). Analysing standard progressive matrices (SPM-LS) with Bayesian item response models. *Journal of Intelligence*, 8, 5. doi:10.3390/jintelligence8010005
- Campos, A. D., Oliveira, H. M., & Soares, A. P. (2020). Syllable effects in beginning and intermediate European-Portuguese readers: Evidence from a sandwich masked go/no-go lexical decision task. *Journal of Child Language*, 1–18. doi:10.1017/S0305000920000537
- Colombo, L., Spinelli, G., & Lupker, S. J. (2020). The impact of consonant–vowel transpositions on masked priming effects in Italian and English. *Quarterly Journal of Experimental Psychology*, 73, 183–198. doi:10.1177/41740720121881986677636838
- Coltheart, M., Davelaar, E., Jonasson, J. T., & Besner, D. (1977). Access to the internal lexicon. In S. Dornic (Ed.), *Attention and Performance VI* (pp. 535–555). Hillsdale, NJ: Erlbaum.
- Coltheart, M., Rastle, K., Perry, C., Langdon, R., & Ziegler, J. (2001). DRC: a dual route cascaded model of visual word recognition and reading

- aloud. *Psychological Review*, 108, 204–256. doi:10.1037/0033-295X.108.1.204
- Comesaña, M., Soares, A. P., Marcet, A., & Perea, M. (2016). On the nature of consonant/vowel differences in letter position coding: Evidence from developing and adult readers. *British Journal of Psychology*, 107, 651–674. doi:10.1111/bjop.12179
- Davis, C. J. (2010). The spatial coding model of visual word identification. *Psychological Review*, 117, 713–758. doi:10.1037/a0019738
- Davis, C. J., & Lupker, S. J. (2006). Masked inhibitory priming in English: Evidence for lexical inhibition. *Journal of Experimental Psychology: Human Perception and Performance*, 32, 668–687. doi:10.1037/0096-1523.32.3.668
- Davis, C., & Forster, K. I. (1994). Masked orthographic priming: The effect of prime-target legibility. *Quarterly Journal of Experimental Psychology*, 47, 673–697. doi:10.1080/14640749408401133
- de Jong, R., Liang, C.-C., & Lauber, E. (1994). Conditional and unconditional automaticity: A dual-process model of effects of spatial stimulus-response correspondence. *Journal of Experimental Psychology: Human Perception and Performance*, 20, 731–750. doi:10.1037//0096-1523.20.4.731
- de Wit, B., & Kinoshita, S. (2014). Relatedness proportion effects in semantic categorization: Reconsidering the automatic spreading activation process. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 40, 1733–1744. doi:10.1037/xlm0000004
- Dehaene, S., Naccache, L., Cohen, L., Le Bihan, D., Mangin, J., Poline, J., & Riviere, D. (2001). Cerebral mechanisms of word masking and unconscious repetition priming. *Nature Neuroscience*, 4, 752–758. doi:10.1038/89551
- Dehaene, S., Naccache, L., Le Clec'h, G., Koechlin, E., Mueller, M., Dehaene-Lambertz, G., ... Le Bihan, D. (1998). Imaging unconscious semantic priming. *Nature*, 395, 597–599. doi:10.1038/26967

- Duchon, A., Perea, M., Sebastián-Gallés, N., Martí, A., & Carreiras, M. (2013). EsPal: One-stop shopping for Spanish word properties. *Behavior Research Methods*, 45, 1246–1258. doi:10.3758/s13428-013-0326-1
- Fernández-López, M., Marcet, A., & Perea, M. (2019). Can response congruency effects be obtained in masked priming lexical decision? *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 45, 1683–1702. doi:10.1037/xlm0000666
- Forster, K. I. (1998). The pros and cons of masked priming. *Journal of Psycholinguistic Research*, 27, 203–233. doi:10.1023/A:1023202116609
- Forster, K. I. (2004). Category size effects revisited: Frequency and masked priming effects in semantic categorization. *Brain and Language*, 90, 276–286. doi:10.1016/S0093-934X(03)00440-1
- Forster, K. I. (2009). The intervenor effect in masked priming: How does masked priming survive across an intervening word?. *Journal of Memory and Language*, 60, 36–49. doi:10.1016/j.jml.2008.08.006
- Forster, K. I. (2013). How many words can we read at once? More intervenor effects in masked priming. *Journal of Memory and Language*, 69, 563–573. doi:10.1016/j.jml.2013.07.004
- Forster, K. I., & Davis, C. (1984). Repetition priming and frequency attenuation in lexical access. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 10, 680–698. doi:10.1037/0278-7393.10.4.680
- Forster, K. I., & Forster, J. C. (2003). DMDX: A Windows display program with millisecond accuracy. *Behavior Research Methods*, 35, 116–124. doi:10.3758/BF03195503
- Forster, K. I., Davis, C., Schoknecht, C., & Carter, R. (1987). Masked priming with graphemically related forms: Repetition or partial activation?. *The Quarterly Journal of Experimental Psychology*, 39, 211–251. doi:10.1080/14640748708401785

- Forster, K. I., Mohan, K., & Hector, J. (2003). The mechanics of masked priming. In S. Kinoshita and S. Lupker (Eds.), *Masked priming: The state of the art* (pp. 3–37). New York: Psychology press.
- Gomez, P., & Perea, M. (2020). Masked identity priming reflects an encoding advantage in developing readers. *Journal of Experimental Child Psychology*, 199, 104911. doi:10.1016/j.jecp.2020.104911
- Gomez, P., Perea, M., & Ratcliff, R. (2013). A diffusion model account of masked versus unmasked priming: Are they qualitatively different? *Journal of Experimental Psychology: Human Perception and Performance*, 39, 1731–1740. doi:10.1037/a0032333
- Grainger, J. (2008). Cracking the orthographic code: An introduction. *Language and Cognitive Processes*, 23, 1–35. doi:10.1080/01690960701578013
- Grainger, J., & Ferrand, L. (1994). Phonology and orthography in visual word recognition: Effects of masked homophone primes. *Journal of Memory and Language*, 33, 218–233. doi:10.1006/jmla.1994.1011
- Grainger, J., & Holcomb, P. J. (2009). Watching the word go by: On the time-course of component processes in visual word recognition. *Language and Linguistics Compass*, 3, 128–156. doi:10.1111/j.1749-818X.2008.00121.x
- Grainger, J., & Jacobs, A. M. (1996). Orthographic processing in visual word recognition: a multiple read-out model. *Psychological Review*, 103, 518–565. doi:10.1037/0033-295X.103.3.518
- Gutierrez-Sigut, E., Vergara-Martínez, M., & Perea, M. (2017). Early use of phonological codes in deaf readers: An ERP study. *Neuropsychologia*, 106, 261–279. doi:10.1016/j.neuropsychologia.2017.10.006
- Heathcote, A., Popiel, S. J., & Mewhort, D. J. (1991). Analysis of response time distributions: An example using the Stroop task. *Psychological Bulletin*, 109, 340–347. doi:10.1037/0033-2909.109.2.340
- Huber, D. E., & O'Reilly, R. C. (2003). Persistence and accommodation in short-term priming and other perceptual paradigms: temporal segregation

- through synaptic depression. *Cognitive Science*, 27, 403–430. doi:10.1207/s15516709cog2703_4
- Jacobs, A. M., & Grainger, J. (1992). Testing a semistochastic variant of the interactive activation model in different word recognition experiments. *Journal of Experimental Psychology: Human Perception and Performance*, 18, 1174–1188. doi:10.1037/0096-1523.18.4.1174
- Jacobs, A. M., Grainger, J., & Ferrand, L. (1995). The incremental priming technique: A method for determining within-condition priming effects. *Perception, & Psychophysics*, 57, 1101–1110. doi:10.3758/BF03208367
- Keuleers, E., & Brysbaert, M. (2010). Wuggy: A multilingual pseudoword generator. *Behavior Research Methods*, 42, 627–633. doi:10.3758/BRM.42.3.627
- Kinoshita, S., & Hunt, L. (2008). RT distribution analysis of category congruence effects with masked primes. *Memory & Cognition*, 36, 1324–1334. doi:10.3758/mc.36.7.1324
- Kinoshita, S., & Norris, D. (2012). Task-dependent masked priming effects in visual word recognition. *Frontiers in Psychology*, 3, 1–12. doi:10.3389/fpsyg.2012.00178
- Ktori, M., Grainger, J., Dufau, S., & Holcomb, P. J. (2012). The “electrophysiological sandwich”: A method for amplifying ERP priming effects. *Psychophysiology*, 49, 1114–1124. doi:10.1111/j.1469-8986.2012.01387.x
- Ktori, M., Kingma, B., Hannagan, T., Holcomb, P. J., & Grainger, J. (2014). On the time-course of adjacent and non-adjacent transposed-letter priming. *Journal of Cognitive Psychology*, 26, 491–505. doi:10.1080/20445911.2014.922092
- Kuznetsova, A., Brockhoff, P. B., & Christensen, R. H. B. (2016). *lmerTest: Tests for random and fixed effects for linear mixed effect models (lmer objects of lme4 package)*. R package Version 2.0 –33. Retrieved from <http://CRAN.Rproject.org/package=lmerTest>

- Lehtonen, M., Monahan, P. J., & Poeppel, D. (2011). Evidence for early morphological decomposition: combining masked priming with magnetoencephalography. *Journal of Cognitive Neuroscience*, 23, 3366–3379. doi:10.1162/jocn_a_00035
- Lenth, R., Singmann, H., Love, J., Buerkner, P., & Herve, M. (2020). emmeans: Estimated marginal means (R package version 1.4.8.). [https:// CRAN.R-project.org/package=emmeans](https://CRAN.R-project.org/package=emmeans)
- Lété, B., & Fayol, M. (2013). Substituted-letter and transposed-letter effects in a masked priming paradigm with French developing readers and dyslexics. *Journal of Experimental Child Psychology*, 114, 47–62. doi:10.1016/j.jecp.2012.09.001
- Loth, S., & Davis, C. J. (2010a, January). *Response congruency effects in masked primed lexical decision*. Poster presented at the Experimental Psychology Society, London Meeting, London, UK.
- Loth, S., & Davis, C. J. (2010b). Testing computational accounts of response congruency in lexical decision. In E. J., Davelaar (Ed.), *Connectionist models of neurocognition and emergent behavior: From theory to applications* (pp. 173–189). Singapore: World Scientific Publishing. doi:10.1142/9789814340359_0012
- Lupker, S. J., & Davis, C. J. (2009). Sandwich priming: A method for overcoming the limitations of masked priming by reducing lexical competitor effects. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 35, 618–639. doi:10.1037/a0015278
- Lupker, S. J., Spinelli, G., & Davis, C. J. (2020a). Masked form priming as a function of letter position: An evaluation of current orthographic coding models. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 46, 2349–2366. doi:10.1037/xlm0000799
- Lupker, S. J., Spinelli, G., & Davis, C. J. (2020b). Is zjudge a better prime for JUDGE than zudge is?: A new evaluation of current orthographic coding

- models. *Journal of Experimental Psychology: Human Perception and Performance*, 46, 1252–1266. doi:10.1037/xhp0000856
- Lupker, S. J., Zhang, Y. J., Perry, J. R., & Davis, C. J. (2015). Superset versus substitution-letter priming: An evaluation of open-bigram models. *Journal of Experimental Psychology: Human Perception and Performance*, 4, 138–151. doi:10.1037/a0038392
- McClelland, J. L., & Rumelhart, D. E. (1981). An interactive activation model of context effects in letter perception: I. An account of basic findings. *Psychological Review*, 88, 375–407. doi:10.1037/0033-295X.88.5.375
- Meade, G., Grainger, J., Midgley, K. J., Holcomb, P. J., & Emmorey, K. (2020). An ERP investigation of orthographic precision in deaf and hearing readers. *Neuropsychologia*, 146, 107542. doi:10.1016/j.neuropsychologia.2020.107542
- Monahan, P. J., Fiorentino, R., & Poeppel, D. (2008). Masked repetition priming using magnetoencephalography. *Brain and Language*, 106, 65–71. doi:10.1016/j.bandl.2008.02.002
- Naccache, L., & Dehaene, S. (2001). Unconscious semantic priming extends to novel unseen stimuli. *Cognition*, 80, 215–229. doi:10.1016/S0010-0277%2800%2900139-6
- Norris, D. (2006). The Bayesian reader: Explaining word recognition as an optimal Bayesian decision process. *Psychological Review*, 113, 327–357. doi:10.1037/0033-295X.113.2.327
- Norris, D. (2009). Putting it all together: A unified account of word recognition and reaction-time distributions. *Psychological Review*, 116, 207–219. doi:10.1037/a0014259
- Norris, D., & Kinoshita, S. (2008). Perception as evidence accumulation and Bayesian inference: Insights from masked priming. *Journal of Experimental Psychology: General*, 137, 434–455. doi:10.1037/a0012799

- Ollman R. T. (1977). Choice reaction time and the problem of distinguishing task effects from strategy effects. In S. Domic (Ed.), *Attention & Performance VI* (pp. 99–113). Hillsdale, NJ: Erlbaum.
- Perea, M., & Rosa, E. (2000). Repetition and form priming interact with neighborhood density at a brief stimulus onset asynchrony. *Psychonomic Bulletin & Review*, 7, 668–677. doi:10.3758/bf03213005
- Perea, M., Abu Mallouh, & Carreiras, M. (2014). Are root letters compulsory for lexical access in Semitic languages? The case of masked form-priming in Arabic. *Cognition*, 132, 491–500. doi:10.1016/j.cognition.2014.05.008
- Perea, M., Fernández, L., & Rosa, E. (1998). El papel del status léxico y de la frecuencia del estímulo-señal en la condición no relacionada con la técnica de presentación enmascarada del estímulo-señal. *Psicológica*, 19, 311–319.
- Perea, M., Gómez, P., & Fraga, I. (2010). Masked nonword repetition effects in yes/no and go/no-go lexical decision: A test of the evidence accumulation and deadline accounts. *Psychonomic Bulletin & Review*, 17, 369–374. doi:10.3758/PBR.17.3.369
- Perea, M., Jiménez, M., & Gomez, P. (2014). A challenging dissociation in masked identity priming with the lexical decision task. *Acta Psychologica*, 148, 130–135. doi:10.1016/j.actpsy.2014.01.014
- Perea, M., Marcet, A., Lozano, M., & Gomez, P. (2018). Is masked priming modulated by memory load? A test of the automaticity of masked identity priming in lexical decision. *Memory & Cognition*, 46, 1127–1135. doi:10.3758/s13421-018-0825-5
- Pollatsek, A., & Treiman, R. (Eds.). (2015). *The Oxford handbook of reading*. Oxford University Press.
- R Core Team (2017). *R: A language and environment for statistical computing*. Vienna: R Foundation for Statistical Computing. Retrieved from <http://www.R-project.org/>

- R Core Team (2020). *R: A language and environment for statistical computing*. Vienna: R Foundation for Statistical Computing. Retrieved from <http://www.R-project.org/>
- Ratcliff, R. (1979). Group reaction time distributions and an analysis of distribution statistics. *Psychological Bulletin*, 86, 446–461. doi:10.1037/0033-2909.86.3.446
- Ratcliff, R., Gomez, P., & McKoon, G. (2004). A diffusion model account of the lexical decision task. *Psychological Review*, 111, 159–182. doi:10.1037/0033-295X.111.1.159
- Ridderinkhof, K. R., van den Wildenberg, W. P. M., Wijnen, J., & Burle, B. (2004). *Response Inhibition in Conflict Tasks Is Revealed in Delta Plots*. In M. I. Posner (Ed.), *Cognitive Neuroscience of Attention* (p. 369–377). The Guilford Press.
- Rouder, J. N., Speckman, P. L., Sun, D., Morey, R. D., & Iverson, G. (2009). Bayesian t tests for accepting and rejecting the null hypothesis. *Psychonomic Bulletin & Review*, 16, 225–237. doi:10.3758/PBR.16.2.225
- Rubinstein, H., Garfield, L., & Millikan, J. A. (1970). Homographic entries in the internal lexicon. *Journal of Verbal Learning and Verbal Behavior*, 9, 487–494. doi:10.1016/S0022-5371(70)80091-3
- Rubinstein, H., Lewis, S. S. & Rubinstein, M. A. (1971a). Homographic entries in the internal lexicon: Effects of systematicity and relative frequency of meanings. *Journal of Verbal Learning and Verbal Behavior*, 10, 57–62. doi:10.1016/S0022-5371(71)80094-4
- Rubinstein, H., Lewis, S. S. & Rubinstein, M. A. (1971b). Evidence for phonemic recoding in visual word recognition. *Journal of Verbal Learning and Verbal Behavior*, 10, 645–657. doi:10.1016/S0022-5371(71)80071-3
- Sereno, J. A. (1991). Graphemic, associative, and syntactic priming effects at a brief stimulus onset asynchrony in lexical decision and naming. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 17, 459–477. doi:10.1037/0278-7393.17.3.459

- Stinchcombe, E. J., Lupker, S. J., & Davis, C. J. (2012). Transposed-letter priming effects with masked subset primes: A re-examination of the “relative position priming constraint”. *Language and Cognitive Processes*, 27, 475–499. doi:10.1080/01690965.2010.550928
- Taikh, A., & Lupker, S. J. (2020). Do visible semantic primes preactivate lexical representations? *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 46, 1533–1569. doi:10.1037/xlm0000825
- Townsend, J. T., & Ashby, F. G. (1983). *Stochastic Modeling of Elementary Psychological Processes*. Cambridge, England: Cambridge University Press.
- Trifonova, I. V., & Adelman, J. S. (2018). The sandwich priming paradigm does not reduce lexical competitor effects. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 44, 1743–1764. doi:10.1037/xlm0000542
- van Heuven, W. J. (2002). Pseudo (2.07) [software]. Retrieved on January 20, 2018 from <http://www.psychology.nottingham.ac.uk/staff/wvh/internal/research/pseudo/index.html>
- van Heuven, W. J., Mandera, P., Keuleers, E., & Brysbaert, M. (2014). SUBTLEX-UK: A new and improved word frequency database for British English. *Quarterly Journal of Experimental Psychology*, 67, 1176–1190. doi:10.1080/17470218.2013.850521
- Yarkoni, T., Balota, D., & Yap, M. (2008). Moving beyond Coltheart’s N: A new measure of orthographic similarity. *Psychonomic Bulletin & Review*, 15, 971–979. doi:10.3758/PBR.15.5.971
- Ziegler, J. C., Bertrand, D., Lété, B., & Grainger, J. (2014). Orthographic and phonological contributions to reading development: Tracking developmental trajectories using masked priming. *Developmental Psychology*, 50, 1026–1036. doi:10.1037/a0035187

Appendix: journal version of papers

Fernández-López, M., Gómez, P., & Perea, M. (2021). Which factors modulate letter position coding in pre-literate children? *Frontiers in Psychology*, 12, 708274. doi:10.3389/fpsyg.2021.708274

Fernández-López, M., Marcet, A., & Perea, M. (2021). Does orthographic processing emerge rapidly after learning a new script? *British Journal of Psychology*, 112, 52–91. doi:10.1111/BJOP.12469

Fernández-López, M., Mirault, J., Grainger, J., & Perea, M. (2021). How resilient is reading to letter rotations? *A parafoveal preview investigation. Journal of Experimental Psychology: Learning, Memory, and Cognition*, 47, 2029–2042. doi:10.1037/xlm0000999

Fernández-López, M., Marcet, A., & Perea, M. (2019). Can response congruency effects be obtained in masked priming lexical decision? *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 45, 1683–1702. doi:10.1037/xlm0000666

Fernández-López, M., Davis, C. J., Perea, M., Marcet, A., & Gómez, P. (in press). Unveiling the boost in the sandwich priming technique. *Quarterly Journal of Experimental Psychology*. doi:10.1177/17470218211055097



Which Factors Modulate Letter Position Coding in Pre-literate Children?

María Fernández-López^{1*}, Pablo Gómez² and Manuel Perea^{1,3}

¹Department of Methodology of Behavioral Sciences and ERI-Lectura, Universitat de València, València, Spain, ²Department of Psychology, California State University, San Bernardino, Palm Desert, CA, United States, ³Center of Research in Cognition, Universidad Antonio de Nebrija, Madrid, Spain

OPEN ACCESS

Edited by:

Tânia Fernandes,
Universidade de Lisboa, Portugal

Reviewed by:

Kenneth R. Paap,
San Francisco State University,
United States
Paulo Ventura,
University of Lisbon, Portugal
Thomas Lachmann,
University of Kaiserslautern, Germany

*Correspondence:

María Fernández-López
maria.fernandez@uv.es

Specialty section:

This article was submitted to
Language Sciences,
a section of the journal
Frontiers in Psychology

Received: 11 May 2021

Accepted: 13 July 2021

Published: 05 August 2021

Citation:

Fernández-López M, Gómez P and
Perea M (2021) Which Factors
Modulate Letter Position Coding in
Pre-literate Children?
Front. Psychol. 12:708274.
doi: 10.3389/fpsyg.2021.708274

One of the central landmarks of learning to read is the emergence of orthographic processing (i.e., the encoding of letter identity and letter order): it constitutes the necessary link between the low-level stages of visual processing and the higher-level processing of words. Regarding the processing of letter position, many experiments have shown worse performance in various tasks for the transposed-letter pair judge-JUDGE than for the orthographic control jupte-JUDGE. Importantly, 4-y.o. pre-literate children also show letter transposition effects in a same-different task: TZ-ZT is more error-prone than TZ-PH. Here, we examined whether this effect with pre-literate children is related to the cognitive and linguistic skills required to learn to read. Specifically, we examined the relation of the transposed-letter in a same-different task with the scores of these children in phonological, alphabetic and metalinguistic awareness, linguistic skills, and basic cognitive processes. To that end, we used a standardized battery to assess the abilities related with early reading acquisition. Results showed that the size of the transposed-letter effect in pre-literate children was strongly associated with the sub-test on basic cognitive processes (i.e., memory and perception) but not with the other sub-tests. Importantly, identifying children who may need a pre-literacy intervention is crucial to minimize eventual reading difficulties. We discuss how this marker can be used as a tool to anticipate reading difficulties in beginning readers.

Keywords: learning to read, orthographic processing, cognitive processing, pre-literate, transposed-letter effect

INTRODUCTION

Whereas language is a unique and sophisticated human ability that emerges naturally in children, reading is a learned skill that needs intensive practice. In fact, reading acquisition is a complex process that involves functional brain changes and requires the correct execution of numerous mental functions (see Maurer et al., 2005). For this reason, children must have adequate perceptual and cognitive skills before the initial steps of reading instruction. Once acquired, reading becomes the most important tool for knowledge acquisition in academic settings and beyond.

In alphabetic scripts, readers can quickly map the visual input into abstract letter representations and, subsequently, into word representations (see Dehaene et al., 2005; Grainger et al., 2008). The emergence of these abstract letter representations would occur during the first 2 years of

reading acquisition (Jackson and Coltheart, 2001). Consistent with this view, using Forster and Davis (1984) masked priming technique, Gomez and Perea (2020) found that, for Grade 2 readers, the identification time of a word like *EDGE* is virtually the same when rapidly preceded by the physically identical prime *EDGE* and when preceded by the nominally (but not physically) identical prime *edge*.

Importantly, the process of visual word recognition requires not only the encoding of the abstract identity of the letters that compose each word but also the encoding of the serial order of the words' letters. If this process was absent, we would not be able to distinguish similarly spelled words like *spot* and *stop*. Notably, in a recent paper with adult readers, Schmitt and Lachmann (2020), demonstrated that when a target has to be identified in a string, processing occurs in serial order (i.e., from left to right) for letter stimuli, but not for strings composed of letters from an unknown alphabet (Cyrillic and Hebrew). At the same time, a considerable wealth of experiments with children and adults have shown that the encoding of letter order is only approximate: Readers often perceive jumbled words (e.g., *JUGDE* or *CHOLocate*) as the original words (see Perea and Lupker, 2003, 2004; Castles et al., 2007; Guerrero and Forster, 2008; Lupker et al., 2008). As serial order processing is a key component of a wide range of psychological processes, from perception to action (Logan, 2021), it is not surprising that the encoding of serial order is also an essential part of reading and literacy. The main goal of the present study is to shed some light on which cognitive factors are associated with pre-readers' ability to encode letter position accurately.

In the context of reading development, Castles et al. (2007) proposed a "lexical tuning" model in which children encode progressively more precisely the letter positions within words. For instance, in a series of masked priming experiments, they found that the prime *dark* was much more effective at activating *DARK* in Grade 3 than in Grade 6 children. The rationale of this model is that, as reading abilities develop, letter position coding becomes more accurate (see Perfetti's, 2017 lexical quality hypothesis, for a similar claim; but see Grainger and Ziegler, 2011, for a different view).¹ Evidence supporting the lexical tuning model has been obtained not only with children of different ages but also with children of the same age: Better readers encode letter position more accurately than the worse readers (see Gómez et al., 2021; Pagán et al., 2021, for evidence with children and see also Andrews and Lo, 2012; Perea et al., 2016, for parallel evidence with adult readers). Furthermore, a poor encoding of serial order may lead to reading difficulties. Friedmann and Gvion (2001) were the first to report that some individuals present problems at encoding letter position, making frequent errors of letter migration within words—reading broad for board. This deficit, which has been termed "letter position" dyslexia, has been found in many different languages,

including English (Kohnen et al., 2012; see Güven and Friedmann, 2019, for a recent review).

Somewhat surprisingly, examining how the encoding of letter order emerges in young readers and whether some preexisting abilities may help encode serial order in pre-readers has been overlooked in the literature. One of the few exceptions is the longitudinal experiment conducted by Duñabeitia et al. (2015). They used a same-different task with two sequentially presented four-letter strings, and children had to decide whether the letter strings were the same or different. The "different" trials were composed of pairs with two transposed letters (transposed-letter pairs; e.g., *rzsk-rszk*) and pairs with two replaced letters (replacement-letter pairs; e.g., *rzsk-rhck*). If letter position coding is flexible, transposed-letter pairs would be perceptually more similar than replacement-letter pairs, thus producing worse performance (e.g., more false positives). The "transposed-letter" effect is the difference in performance between these two conditions. Duñabeitia et al. (2015) tested the children three times: (1) in their year before preschool ($M = 4.24$ years) (2) in their preschool year ($M = 5.21$ years), and (3) in the first year of primary school ($M = 6.32$ years). They only found a transposed-letter effect when the children had learned to read (first-grade children; more error responses for transposed [42.9%] vs. replaced-letter pairs [30.6%]). Duñabeitia et al. (2015) concluded that "position uncertainty emerges as a consequence of literacy training" (p. 549).

An interpretive issue in the Duñabeitia et al. (2015) experiment is that the pre-readers performed very poorly in the same-different task and the sensitivity index, d' , was close to zero for both for replaced and transposed conditions (Perea et al., 2016). This pattern suggests that their version of the same-different task was too difficult for the pre-readers (i.e., the working memory load probably exceeded the children's capacity; see Riggs et al., 2006); thus, one cannot make any inferences on these data. To draw firm conclusions on the encoding of the serial position of letters in pre-readers, Perea et al. (2016) simplified some elements of Duñabeitia et al.'s (2015) same-different task: (1) they used two-letter string pairs instead of four-letter string pairs (2) the pairs were presented simultaneously instead of sequentially, and (3) the responses were done verbally (i.e., saying "same" vs. "different") instead of manually (pressing one of two buttons; see **Figure 1**). Along with "same" pairs (TZ-TZ), Perea et al. (2016) included the following "different" pairs: transposed-letter pairs (TZ-ZT), one-letter replacement pairs (TZ-PZ), and two-letter replacement pairs (TZ-PH). They found a sizeable transposed-letter effect in 4-years-old children (i.e., pre-readers). Specifically, the number of false positives (i.e., "same" responses) was greater to transposed-letter strings (TZ-ZT) than to 1 or 2 replacement-letter strings (TZ-PZ; TZ-PH). Perea et al. (2016) concluded that this pattern reflected a noisy perception of location order, common to all visual objects (see Gomez et al., 2008), rather than an effect that emerges with literacy. Notably, while not analyzed in their paper, shortly after conducting their experiment, Perea et al. (2016) collected the scores of these children in a battery of abilities related to early reading acquisition in Spanish (BIL battery; Sellés et al., 2008).

¹In fairness to Grainger and Ziegler (2011), their dual-route model of visual word recognition focuses on the initial steps of learning to read. The serial letter encoding, which involves precise letter position encoding, would emerge first in reading development (phonological route). Later in development, the parallel encoding of the word's letters (orthographic route) would make letter position coding coarser.

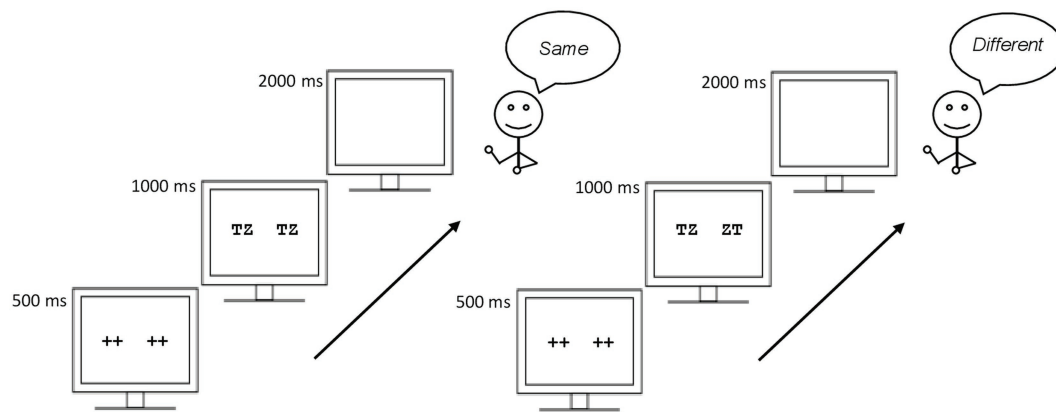


FIGURE 1 | Depiction of the same-different task used in the Perea et al. (2016) study.

In the present study, we aim to take a step forward by exploring the potential precursors of letter position coding in pre-readers. To that end, we examined the relationship between the ability of pre-literate children to encode accurately the order of letters—taken from the Perea et al. (2016) experiment—with the five sub-tests related to reading readiness and subsequent reading success from the BIL battery: phonological and alphabetic awareness, metalinguistic knowledge, linguistic skills, and basic cognitive processes. The examination of this issue is important not only at a theoretical level but also at a practical level. Before learning to read, children must have acquired some perceptual, cognitive, and linguistic skills. Defining the early precursors of precise coding of letter position will shed light on the roots of the processing of serial order when reading letters in words. These analyses would allow us to identify children who may present some deficit (e.g., some mild forms of letter position dyslexia) and start intervening as soon as possible, preventing future reading difficulties.

Thus, in the present study, we examined the relationship between the sensitivity of the readers to distinguish transposed-letter pairs from identity pairs (e.g., TZ-ZT vs. TZ-TZ) and the scores of pre-readers ($M = 4.5$ years old) in phonological and alphabetic awareness, metalinguistic knowledge, linguistic skills, and basic cognitive processes (visual perception and sequential auditory memory) in the BIL battery (Sellés et al., 2008). We focused on transposed-letter pairs, as the mechanisms employed to discriminate TZ-ZT from ZT-ZT are based exclusively on letter order. The predictions are clear. In adult readers, basic cognitive processes, such as spatial and visual attention, have been assumed to play a key role in encoding letter position (see McCann et al., 1992; Gomez et al., 2008). If this generalizes to pre-readers, we expect a positive relationship between the abilities at discriminating TZ-ZT and the scores in these basic cognitive processes. This outcome would imply that educators could use this simple same-different task with a transposed-letter pairs to predict reading readiness before starting with the reading instruction. Furthermore, it may also operate as an incentive to design other tasks for pre-readers on perceptive and executive skills to prevent—or at least minimize—potential difficulties at

locating letters within words during learning to read. In addition, we expect no relation between linguistic factors (i.e., phonological and alphabetic awareness, metalinguistic knowledge, and linguistic skills) and the sensitivity at distinguishing TZ-ZT in pre-readers—at the time of the experiment, the children did not know the consonant names.

MATERIALS AND METHODS

Participants

They were the 20 preschoolers ($M = 4.54$ years; $SD = 3.6$; 7 girls) from a private school of Valencia (Spain). All of them were native speakers of Spanish with no learning developmental problems. An informed consent from their parents was obtained before running the experiment, and the study was approved by the Experimental Research Ethics Committee of the University of Valencia. At the time of testing, the preschoolers were starting to learn the vowels but they did not know the name or sound of the consonant letters (as confirmed by results of the BIL battery).

Procedure

The experiment took place individually in a quiet room within the school premises. DMDX software (Forster and Forster, 2003) was employed for stimulus presentation and recording of the responses. A depiction of the procedure in the same-different task can be found in **Figure 1**. Accuracy was stressed in the instructions. Ten practice trials preceded the 64 experimental trials. Moreover, the children were assessed with a battery of abilities related to early reading acquisition in Spanish (BIL battery; Sellés et al., 2008).

Materials

For the same-different task, the stimuli were 64 pairs of consonant strings made of two consonants. There were 16 trials in each of the conditions: (1) same pairs (TZ-TZ) (2) transposed-letter pairs (TZ-ZT) (3) one-letter replacement

pairs (TZ-PZ), and (4) two-letter replacement pairs (TZ-PH). Four counterbalanced lists were created in a Latin square manner, so that each stimulus was rotated across the different conditions. The presentation of the items was randomized for each participant.

To assess the abilities related to early reading acquisition, we employed the BIL battery (Sellés et al., 2008). This battery comprises five sub-tests: phonological awareness, alphabetic awareness, metalinguistic knowledge, linguistic skills, and basic cognitive processes—we obtained a score from each sub-test. For the goals of the present study, we focused on the sub-test measuring basic cognitive processes. This sub-test explores a series of cognitive processes that take place when we face reading: (1) attention, which leads the mind to concentrate on specific stimuli; (2) sensation (i.e., detection and differentiation of sensory information); and (3) perception, which integrates sensory experiences and interprets them for recognition and identification (i.e., giving meaning to what has been selected and picked up at the attentional and sensory level), relying on the patterns stored in the (4) memory. To that end, the sub-test assesses the child's sequential auditory memory and the ability to visually discriminate between similar letters and symbols (the child had to circle the symbols that were the same as a target; see Sellés et al., 2008, for a depiction of the other sub-tests).

RESULTS

To test whether better pre-reading skills (as measured by the BIL battery) were associated with better performance at differentiating between same and transposed-letter pairs in the same-different task, we conducted frequentist and Bayesian correlation analyses with JASP (Faulkenberry et al., 2020). To compute the Bayes factors, we used the default Cauchy distribution (centered around 0 and with a width parameter $\delta = 0.707$; see Rouder et al., 2009; Wagenmakers et al., 2017, 2018, for discussion). Specifically, we examined the relation between d' (a measure of sensitivity obtained from the accuracy data of Perea et al., 2016) and the percentile scores in sub-tests of the BIL battery—of note, these findings were virtually the same if we had employed the raw scores from the sub-scales. For the computation of d' , we used the hit rate for same trials and the false alarm rate for the transposed-letter trials (TZ-TZ vs. TZ-ZT)—in signal detection theory, chance-level performance [$d' = 0$ or no sensitivity] occurs when the hit rate for the identical items is the equal to the false alarm rate for the different items. Of note, mean accuracy for same trials in the Perea et al. (2016) was 0.83; for different trials, it was 0.33 for transposed-letter strings and 0.68 and 0.88 for one-letter and two-letter replacement strings, respectively.

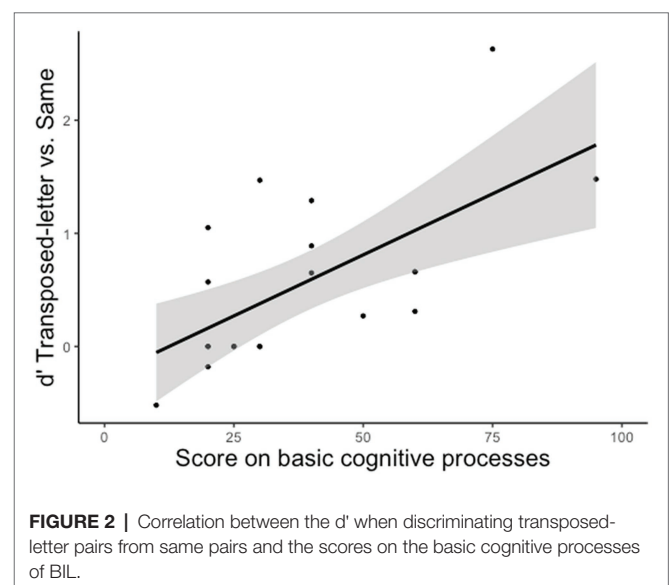
Results of the correlational analyses in the present study showed that those children who better differ transposed letter from “same” pairs (TZ-ZT vs. TZ-TZ) had the higher scores in the sub-test on basic cognitive processes ($r = 0.634$, $p = 0.003$; see Figure 2). Indeed, the alternative hypothesis was 18.6 ($BF_{10} = 18.559$) times more likely than the null hypothesis

with the present data (see Jeffreys, 1961, for interpretation of Bayes factors). In addition, there were no signs of a relationship between the children's performance differentiating between same and transposed-letter pairs and the other (linguistic) sub-tests (all $ps > 0.24$, $BF_{10} < 0.528$).

For completeness, we explored the relationship between performance in the replacement-letter conditions and the BIL battery; to this end, we computed separately d' s for same vs. one-letter replacement trials (TZ-TZ vs. TZ-PZ) and for same vs. two-letter replacement trials (TZ-TZ vs. TZ-PH), and then calculated the correlations between these two d' s and the sub-test of the BIL battery. None of these correlations produced evidence in favor of a relationship (all $ps > 0.147$; all $BF_{10} < 0.478$).

DISCUSSION

Identifying the cognitive precursors of reading is vital to determine those children who are ready to start learning to read and those who still need some cognitive maturation or some early intervention. This would prevent later reading difficulties and disorders and the frustration and psychological discomfort that such problems usually entail. With this matter in mind, in the present study, we scrutinized the roots of the mechanisms underlying the encoding of letter position in strings (i.e., one of the critical factors of efficient reading; see Castles et al., 2007; Logan, 2021). Specifically, we examined the relationship between the capability of pre-literate to differentiate between transposed-letter pairs and identity pairs (e.g., TZ-ZT vs. TZ-TZ) and these children's scores in basic cognitive processes. Results showed that the pre-literate children who best differentiated between TZ-ZT and TZ-TZ in a same-different task were those with higher scores on basic cognitive processes (see Figure 2). Notably, the sub-test of basic cognitive processes was not generically associated with sensitivity in the same-different task (i.e., it



was not related to performance for replacement-letter trials); instead, it is uniquely associated with accuracy in letter position coding. Thus, at the theoretical level, this outcome reflects that basic cognitive skills shape the ability to encode serial order in letter strings (e.g., a smaller value of the parameter responsible for perceptual uncertainty in models of letter position coding; see Gomez et al., 2008; Davis, 2010). Furthermore, at an applied/educational level, our findings imply that a simple same-different task can be used to assess reading readiness: the better the performance in this task, the better the encoding of letter order, diminishing the chances of letter position dyslexia.

In addition, our findings suggest that the preparing-to-reading arises early in development with some non-specialized processes that would be recruited and adjusted to guide the subsequent functional reading progress (see Lachmann and van Leeuwen, 2014). Further support to this idea can be found in the study of Saygin et al. (2016). They found that the cortical location of the visual word form area (i.e., the brain region specialized for letter string; Dehaene and Cohen, 2011) at age 8 (when children read) can be predicted by the distinctive connectivity of the same region at age 5 (pre-literates). Taken together, these studies emphasize that early detection of deficiencies in the visual analysis of the input is crucial to prevent later reading difficulties (Friedmann and Gvion, 2001; Shetreet and Friedmann, 2011). This is consistent with the assumption that children with reading difficulties have a general impairment in domains other than linguistic (e.g., an impairment in multisensory integration; see Gori and Facoetti, 2014; Lachmann and van Leeuwen, 2014). Therefore, there are possibly many (complementary) ways to test whether pre-readers are prepared to starting reading learning (e.g., the same-different task or the “avatar task”; see Perea et al., 2014); this would be a valuable endeavor for the future studies.

We acknowledge that the present study comes with some limitations. Firstly, because of the correction criteria of the BIL test, it was not possible to obtain separate scores for the tasks that make up the basic cognitive processes sub-test (sequential auditory memory and visual discrimination). Furthermore, although the processes assessed in the basic cognitive sub-test were not linguistic in nature, the stimuli contained symbols, letters, and words, thus, making it difficult to clearly disentangle basic cognitive processes and verbal processes. Future tests should be more specific to characterize all possible aspects that shape the cognitive processes of pre-literates. In addition, it would have been desirable to have obtained further data from the same children once they started reading learning. These data would have allowed us to test whether the findings in the same-different task with transposed letters in pre-literate children were a good predictor of letter position coding once children acquired knowledge about letters. Furthermore, these longitudinal data would have also allowed us to examine the interplay between the emergence of orthographic processing during learning and the scores in cognitive and linguistic processes. Indeed, once the children start to read, other elements would begin playing a significant

role, such as alphabetic knowledge or phonologic awareness (see Dehaene et al., 2015).

A complementary strategy for the future research would be to run parallel longitudinal same-different experiments on serial order using to-be-learned letters vs. unknown letters (e.g., letters from another alphabet). The data pattern should be similar for the pre-readers for both types of stimuli, but one would expect differences when the children learn to read. Critically, these differences could be considered as markers of the emergence of orthographic processing (see Grainger, 2018). While this approach is ideal on an *a priori* basis, it suffers from various methodological issues. One would need to design a feasible task for children of different tasks that minimizes both ground and ceiling effects. However, it is challenging to create a task achievable for pre-literates and complicated enough to draw differences among developing readers. For instance, deciding whether two four-letter strings are the same is extremely challenging for pre-literates, whereas deciding whether two-letter strings are the same may be too easy for developing readers (see Perea et al., 2016, for discussion). As a result, it is very difficult to experimentally study the emergence and development of orthographic processes in pre-readers. To further complicate matters, there are also other potential limitations, such as the lack of control for prior letter knowledge and other linguistic elements in pre-readers, or that the duration of experiments for pre-readers would need to be quite short to keep them attentive. An alternative is to design laboratory analogs of children's reading acquisition that consists of training adults to read a novel script (see Fernández-López et al., 2021; see also Chetail, 2017; Taylor et al., 2017). This approximation is not as ecological as one would desire (see Maurer et al., 2010; Taylor et al., 2017), but it definitively increases the control on the process of acquiring the novel orthography.

In sum, the early identification of potential problems that may slow down reading development is of fundamental importance for psychologists and educators. In the present study, we found that those pre-readers who performed better in basic cognitive processes tended to be those who would encode more accurately letter position in a simple same-different task. This finding highlights that learning to read should not be based solely on letter knowledge and phonological decoding. We also need to consider that learning to read is built on a basic cognitive foundation, probably related to multisensory integration based on visual attention.

DATA AVAILABILITY STATEMENT

Publicly available datasets were analyzed in this study. This data can be found at: https://osf.io/hz7m2/?view_only=2073b5df420e4d5f95c627ec4ee6e81b.

ETHICS STATEMENT

The studies involving human participants were reviewed and approved by Comité de Ética de Investigación en Humanos (CEIH).

Written informed consent to participate in this study was provided by the participants' legal guardian/next of kin.

AUTHOR CONTRIBUTIONS

MP, PG and MF-L contributed to conception, design of the study, and performed the statistical analysis. PG organized the database. MF-L wrote the first draft of the manuscript. MF-L and MP wrote sections of the manuscript. All authors

contributed to manuscript revision, read, and approved the submitted version.

FUNDING

This study was supported by the Spanish Ministry of Science and Innovation (grants PRE2018-083922 and PSI2017-86210-P) and the Department of Innovation, Universities, Science, and Digital Society of the Valencian Government (GV/2020/074).

REFERENCES

- Andrews, S., and Lo, S. (2012). Not all skilled readers have cracked the code: Individual differences in masked form priming. *J. Exp. Psychol. Learn. Mem. Cogn.* 38, 152–163. doi: 10.1037/a0024953
- Castles, A., Davis, C., Cavalot, P., and Forster, K. (2007). Tracking the acquisition of orthographic skills in developing readers: masked priming effects. *J. Exp. Child Psychol.* 97, 165–182. doi: 10.1016/j.jecp.2007.01.006
- Chetail, F. (2017). What do we do with what we learn? statistical learning of orthographic regularities impacts written word processing. *Cognition* 163, 103–120. doi: 10.1016/j.cognition.2017.02.015
- Davis, C. J. (2010). The spatial coding model of visual word identification. *Psychol. Rev.* 117, 713–758. doi: 10.1037/a0019738
- Dehaene, S., and Cohen, L. (2011). The unique role of the visual word form area in reading. *Trends Cogn. Sci.* 15, 254–262. doi: 10.1016/j.tics.2011.04.003
- Dehaene, S., Cohen, L., Morais, J., and Kolinsky, R. (2015). Illiterate to literate: behavioural and cerebral changes induced by reading acquisition. *Nat. Rev. Neurosci.* 16, 234–244. doi: 10.1038/nrn3924
- Dehaene, S., Cohen, L., Sigman, M., and Vinckier, F. (2005). The neural code for written words: a proposal. *Trends Cogn. Sci.* 9, 335–341. doi: 10.1016/j.tics.2005.05.004
- Duñabeitia, J. A., Lallier, M., Paz-Alonso, P. M., and Carreiras, M. (2015). The impact of literacy on position uncertainty. *Psychol. Sci.* 26, 548–550. doi: 10.1177/0956797615569890
- Faulkenberry, T. J., Ly, A., and Wagenmakers, E. J. (2020). Bayesian inference in numerical cognition: a tutorial using JASP. *J. Numer. Cogn.* 6, 231–259. doi: 10.5964/jnc.v6i2.288
- Fernández-López, M., Marcet, A., and Perea, M. (2021). Does orthographic processing emerge rapidly after learning a new script? *Br. J. Psychol.* 112, 52–91. doi: 10.1111/bjop.12469
- Forster, K. I., and Davis, C. (1984). Repetition priming and frequency attenuation in lexical access. *J. Exp. Psychol. Learn. Mem. Cogn.* 10, 680–698. doi: 10.1037/0278-7393.10.4.680
- Forster, K. I., and Forster, J. C. (2003). DMDX: a windows display program with millisecond accuracy. *Behav. Res. Methods Inst. Comp.* 35, 116–124. doi: 10.3758/BF03195503
- Friedmann, N., and Gvion, A. (2001). Letter position dyslexia. *Cogn. Neuropsychol.* 18, 673–696. doi: 10.1080/02643290143000051
- Gómez, P., Marcet, A., and Perea, M. (2021). Are better young readers more likely to confuse their mother with their mother? *Q. J. Exp. Psychol.* doi: 10.1177/17470218211012960 (in press).
- Gomez, P., and Perea, M. (2020). Masked identity priming reflects an encoding advantage in developing readers. *J. Exp. Child Psychol.* 199:104911. doi: 10.1016/j.jecp.2020.104911
- Gomez, P., Ratcliff, R., and Perea, M. (2008). The overlap model: A model of letter position coding. *Psychol. Rev.* 115, 577–600. doi: 10.1037/a0012667
- Gori, S., and Facoetti, A. (2014). Perceptual learning as a possible new approach for remediation and prevention of developmental dyslexia. *Vis. Res.* 99, 78–87. doi: 10.1016/j.visres.2013.11.011
- Grainger, J. (2018). Orthographic processing: a 'mid-level' vision of reading: The 44th Sir Frederic Bartlett Lecture. *Q. J. Exp. Psychol.* 71, 335–359. doi: 10.1080/17470218.2017.1314515
- Grainger, J., Rey, A., and Dufau, S. (2008). Letter perception: from pixels to pandemonium. *Trends Cogn. Sci.* 12, 381–387. doi: 10.1016/j.tics.2008.06.006
- Grainger, J., and Ziegler, J. C. (2011). A dual-route approach to orthographic processing. *Front. Psychol.* 2:54. doi: 10.3389/fpsyg.2011.00054
- Guerrera, C., and Forster, K. (2008). Masked form priming with extreme transposition. *Lang. Cogn. Process.* 23, 117–142. doi: 10.1080/01690960701579722
- Güven, S., and Friedmann, N. (2019). Developmental letter position dyslexia in Turkish, a morphologically rich and orthographically transparent language. *Front. Psychol.* 10:2401. doi: 10.3389/fpsyg.2019.02401
- Jackson, N., and Coltheart, M. (2001). *Routes to Reading Success and Failure*. Hove, UK: Psychology Press.
- Jeffreys, H. (1961). *Theory of Probability*. 3rd Edn. Oxford: Oxford University Press, Clarendon Press.
- Kohnen, S., Nickels, L., Castles, A., Friedmann, N., and McArthur, G. (2012). When 'slime' becomes 'smile': developmental letter position dyslexia in English. *Neuropsychologia* 50, 3681–3692. doi: 10.1016/j.neuropsychologia.2012.07.016
- Lachmann, T., and van Leeuwen, C. (2014). Reading as functional coordination: not recycling but a novel synthesis. *Front. Psychol.* 5:1046. doi: 10.3389/fpsyg.2014.01046
- Logan, G. D. (2021). Serial order in perception, memory, and action. *Psychol. Rev.* 128, 1–44. doi: 10.1037/rev0000253
- Lupker, S. J., Perea, M., and Davis, C. J. (2008). Transposed-letter effects: consonants, vowels and letter frequency. *Lang. Cogn. Process.* 23, 93–116. doi: 10.1080/01690960701579714
- Maurer, U., Blau, V. C., Yoncheva, Y. N., and McCandliss, B. D. (2010). Development of visual expertise for reading: rapid emergence of visual familiarity for an artificial script. *Dev. Neuropsychol.* 35, 404–422. doi: 10.1080/87565641.2010.480916
- Maurer, U., Brem, S., Bucher, K., and Brandeis, D. (2005). Emerging neurophysiological specialization for letter strings. *J. Cogn. Neurosci.* 17, 1532–1552. doi: 10.1162/089892905774597218
- McCann, R. S., Folk, C. L., and Johnston, J. C. (1992). The role of spatial attention in visual word processing. *J. Exp. Psychol. Hum. Percept. Perform.* 18, 1015–1029. doi: 10.1037/0096-1523.18.4.1015
- Pagán, A., Blythe, H. I., and Liversedge, S. P. (2021). The influence of children's reading ability on initial letter position encoding during a reading-like task. *J. Exp. Psychol. Learn. Mem. Cogn.* doi: 10.1037/xlm0000989 [Epub ahead of print]
- Perea, M., Jiménez, M., and Gomez, P. (2016). Does location uncertainty in letter position coding emerge because of literacy training? *J. Exp. Psychol. Learn. Mem. Cogn.* 42, 996–1001. doi: 10.1037/xlm0000208
- Perea, M., Jiménez, M., Suárez-Coalla, P., Fernández, N., Viña, C., and Cueto, F. (2014). Ability for voice recognition is a marker for dyslexia in children. *Exp. Psychol.* 61, 480–487. doi: 10.1027/1618-3169/a000265
- Perea, M., and Lupker, S. J. (2003). Does judge activate COURT? Transposed-letter similarity effects in masked associative priming. *Mem. Cogn.* 31, 829–841. doi: 10.3758/BF03196438
- Perea, M., and Lupker, S. J. (2004). Can CANISO activate CASINO? transposed-letter similarity effects with nonadjacent letter positions. *J. Mem. Lang.* 51, 231–246. doi: 10.1016/j.jml.2004.05.005
- Perea, M., Marcet, A., and Gomez, P. (2016). How do Scrabble players encode letter position during reading? *Psicothema* 28, 7–12. doi: 10.7334/psicothema2015.167
- Perfetti, C. A. (2017). "Lexical quality revisited," in *Developmental Perspectives in Written Language and Literacy: In Honor of Ludo Verhoeven*. eds.

- E. Segers and P. van den Broek (Amsterdam, Netherlands: John Benjamins), 51–68).
- Riggs, K. J., McTaggart, J., Simpson, A., and Freeman, R. P. (2006). Changes in the capacity of visual working memory in 5- to 10-year-olds. *J. Exp. Child Psychol.* 95, 18–26. doi: 10.1016/j.jecp.2006.03.009
- Rouder, J. N., Speckman, P. L., Sun, D., Morey, R. D., and Iverson, G. (2009). Bayesian t tests for accepting and rejecting the null hypothesis. *Psychon. Bull. Rev.* 16, 225–237. doi: 10.3758/PBR.16.2.225
- Saygin, Z. M., Osher, D. E., Norton, E. S., Youssoufian, D. A., Beach, S. D., Feather, J., et al. (2016). Connectivity precedes function in the development of the visual word form area. *Nat. Neurosci.* 19, 1250–1255. doi: 10.1038/nn.4354
- Schmitt, A., and Lachmann, T. (2020). Skilled readers show different serial-position effects for letter versus non-letter target detection in mixed-material strings. *Acta Psychol.* 204:103025. doi: 10.1016/j.actpsy.2020.103025
- Sellés, P., Martínez, T., Vidal-Abarca, E., and Gilabert, R. (2008). *BIL 3–6. Bateria de Inicio a la Lectura [Battery to Assess the Abilities Related with Early Reading Acquisition]*. Madrid, Spain: Instituto de Ciencias de la Educación.
- Shetreet, E., and Friedmann, N. (2011). Induced letter migrations between words and what they reveal about the orthographic-visual analyzer. *Neuropsychologia* 49, 339–351. doi: 10.1016/j.neuropsychologia.2010.11.026
- Taylor, J. S. H., Davis, M. H., and Rastle, K. (2017). Comparing and validating methods of reading instruction using behavioural and neural findings in an artificial orthography. *J. Exp. Psychol. Gen.* 146, 826–858. doi: 10.1037/xge0000301
- Wagenmakers, E. J., Love, J., Marsman, M., Jamil, T., Ly, A., Verhagen, J., et al. (2018). Bayesian inference for psychology. Part II: Example applications. *JASP. Psychon. Bull. Rev.* 25, 58–76. doi: 10.3758/s13423-017-1323-7
- Wagenmakers, E.-J., Marsman, M., Jamil, T., Ly, A., Verhagen, J., Love, J., et al. (2017). Bayesian inference for psychology. Part I: Theoretical advantages and practical ramifications. *Psychon. Bull. Rev.* 25, 35–57. doi: 10.3758/s13423-017-1343-3

Conflict of Interest: The authors declare that the research was conducted in the absence of any commercial or financial relationships that could be construed as a potential conflict of interest.

Publisher's Note: All claims expressed in this article are solely those of the authors and do not necessarily represent those of their affiliated organizations, or those of the publisher, the editors and the reviewers. Any product that may be evaluated in this article, or claim that may be made by its manufacturer, is not guaranteed or endorsed by the publisher.

Copyright © 2021 Fernández-López, Gómez and Perea. This is an open-access article distributed under the terms of the Creative Commons Attribution License (CC BY). The use, distribution or reproduction in other forums is permitted, provided the original author(s) and the copyright owner(s) are credited and that the original publication in this journal is cited, in accordance with accepted academic practice. No use, distribution or reproduction is permitted which does not comply with these terms.

Registered report

Does orthographic processing emerge rapidly after learning a new script?

María Fernández-López¹ , Ana Marcet¹  and
Manuel Perea^{1,2,3*} 

¹Universitat de València, Spain

²Basque Center on Brain, Cognition, and Language, Donostia, Spain

³Universidad Nebrija, Spain

Orthographic processing is characterized by location-invariant and location-specific processing (Grainger, 2018): (1) strings of letters are more vulnerable to transposition effects than the strings of symbols in same-different tasks (location-invariant processing); and (2) strings of letters, but not strings of symbols, show an initial position advantage in target-in-string identification tasks (location-specific processing). To examine the emergence of these two markers of orthographic processing, we conducted a same-different task and a target-in-string identification task with two unfamiliar scripts (pre-training experiments). Across six training sessions, participants learned to fluently read and write one of these scripts. The post-training experiments were parallel to the pre-training experiments. Results showed that the magnitude of the transposed-letter effect in the same-different task and the serial function in the target-in-string identification tasks were remarkably similar for the trained and untrained scripts. Thus, location-invariant and location-specific processing does not emerge rapidly after learning a new script; instead, they may require thorough experience with specific orthographic structures.

Reading is an acquired skill that involves some functional brain changes and requires, in alphabetic scripts, associating the letters that compose each word with their appropriate speech sounds. A common assumption in the literature is that, for a mature word recognition system, the process of identifying words comprises a series of stages that map the visual input onto abstract letter representations and, subsequently, onto whole-word representations (see Dehaene, Cohen, Sigman, & Vinckier, 2005; Grainger, 2008; but see Price & Devlin, 2011, for an alternative account). The processing of orthographic representations connects the low-level stages of visual processing to the higher-level linguistic processing of words. These orthographic representations contain information about the identity and order of each of the word's constituent letters, thus allowing readers to distinguish similarly spelled words like *cure* and *core*, which differ in the identity of just one letter, or words like *slat* and *salt*, which differ in the order of two of the letters (see Grainger, 2018, for review). Indeed, the question of how the brain encodes the identity and order of the letters that constitute each word is a central issue for all leading

*Correspondence should be addressed to Manuel Perea, Universitat de València, Av Blasco Ibañez, 21, 46010 Valencia, Spain (email: mperea@uv.es).

models of visual word recognition (e.g., Spatial Coding model: Davis, 2010; Overlap model: Gomez, Ratcliff, & Perea, 2008; Open Bigram model: Grainger & Van Heuven, 2004; Bayesian Reader model: Norris, Kinoshita, & van Casteren, 2010; SERIOL model: Whitney, 2001).

In the present experiments, we examined the emergence of two fundamental markers of orthographic processing after learning a new script: 1) location-invariant processing and 2) location-specific processing (see Grainger, 2018). For simplicity, the above-cited models focus on an extant perspective (i.e., they assume a fully developed word recognition system frozen in time) rather than on a developmental perspective. (We defer a discussion of the research focused on the development rather than the emergence of orthographic representations [e.g., Castles, Davis, Cavalot, & Forster, 2007; Grainger, Lété, Bertrand, Dufau, & Ziegler, 2012; Marinus, Kezilas, Kohnen, Robidoux, & Castles, 2018] until the Discussion section.) We first describe in some depth how location-invariant and location-specific processing differs between letters and other visual objects (e.g., symbols, unknown letters). Then, we describe how acquiring a new script may affect both phenomena. Finally, we offer a rationale for the two experiments proposed in the current paper.

Location-invariant processing refers to the mechanisms responsible for encoding the ‘relative positions of a set of object identities’ (Grainger, 2018, p. 345) (i.e., the encoding of the order of visual objects [letters] in a string composed of several objects [a word]). This has often been examined with the same-different matching task (see Krueger, 1978; Ratcliff, 1981, for early evidence), as it allows researchers to compare the processing of letters vs. the processing of other types of visual objects. In this task, participants have to decide if two strings of characters presented subsequently are the same or not (see Figure 1). The most studied phenomenon of location-invariant processing is the transposed-letter effect (henceforth, TL effect; see Grainger, 2018, for review). The TL effect refers to the insensitivity of readers to the position of letters compared to the identity of the same letters: ‘no’ responses to the transposed-letter pair FGJM-FJGM (the underline is to emphasize the manipulation) in a same-different matching task are slower and more error prone than the responses to the replacement-letter control FGJM-FPCM. These effects also occur with strings composed of symbols or unknown letters (e.g., £\$?@-£\$?@ is slower and more error prone than £\$?@-£#<@), which suggests that there is some positional noise in the representations of visual objects in a string (see Gomez *et al.*, 2008; Norris & Kinoshita, 2012). But the critical finding is that transposition effects are substantially larger for strings of letters than for strings of other visual objects (e.g., numbers, symbols, pseudoletters; Duñabeitia, Dimitropoulou, Grainger, Hernández, & Carreiras, 2012; Massol, Duñabeitia, Carreiras, & Grainger, 2013; see also García-Orza, Perea, & Muñoz, 2010; Muñoz, Perea, García-Orza, & Barber, 2012). To explain the greater transposition effect for strings of letters than for strings of other visual stimuli, Grainger (2018; see also Marcet, Perea, Baciero, & Gomez, 2019; Massol *et al.*, 2013) suggested that, on top of positional noise, there is an orthographic-specific mechanism used to encode location-invariant letter-in-word order. Critically, this orthographic-specific mechanism has been posited to emerge with literacy acquisition (Dandurand, Grainger, & Dufau, 2010; Duñabeitia, Lallier, Paz-Alonso, & Carreiras, 2015; Duñabeitia, Orihuela, & Carreiras, 2014). Therefore, when learning a new script, the emergence of location-invariant orthographic processes would produce an increase of letter transposition effects in a same-different task.

Location-specific processing of visual information refers to the parallel processing of the position of characters (e.g., letters) within one object (e.g., a word). This type of

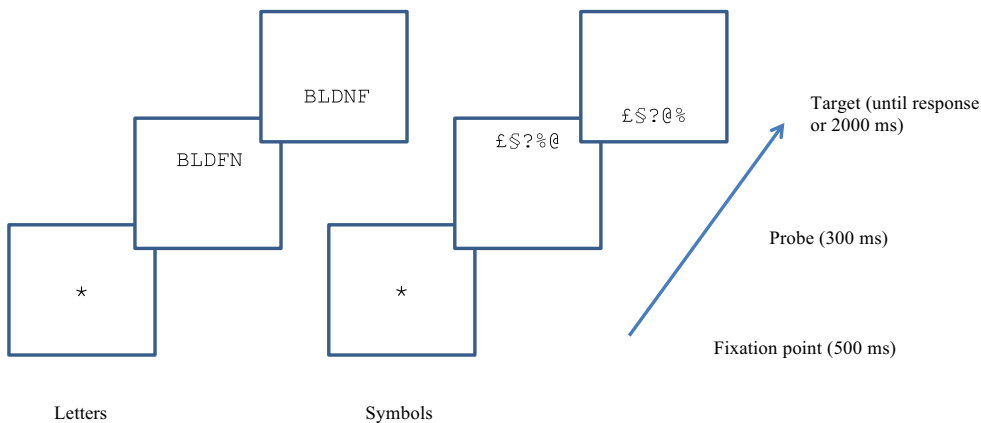


Figure 1. Depiction of the same-different task. [Colour figure can be viewed at wileyonlinelibrary.com]

processing has usually been examined with a post-cued partial report target-in-string identification task (henceforth, TSI task), based on the two-alternative-forced-choice task first introduced by Reicher (1969) and Wheeler (1970). In the typical set-up of the TSI task applied to location-specific processing (e.g., Tydgate & Grainger, 2009), a string of five characters (e.g., letters: FGJGM; symbols: £\$?%@) is presented briefly while the participant is looking at the middle of the string. The string is subsequently followed by a pattern mask with a cue indicating one of the positions in the string (see Figure 2 for illustration). The participants' task is to choose, from the two alternatives, the one that matches the identity of the character at the cued location. When presented with strings of symbols, adult readers typically show a Λ -shape function (i.e., an advantage of the middle, fixated position) (Tydgate & Grainger, 2009; see also Chanceaux & Grainger, 2012; Chanceaux, Mathôt, & Grainger, 2014; Grainger, Tydgate, & Issele, 2010; Scaltritti, Dufau, & Grainger, 2018; Vejnović & Zdravković, 2015). In contrast, for letter strings, adult readers typically show a W-shape serial position function of accuracy (see Tydgate & Grainger, 2009). That is, there is an advantage not only for the fixated, middle letter, but also of the exterior letters – primarily the initial letter. The dissociation in serial position function for strings of symbols vs. letters has also been obtained with developing readers (see Ziegler, Pech-Georgel, Dufau, & Grainger, 2010). To explain this pattern, Tydgate and Grainger (2009) proposed the Modified Receptive Field (MRF) theory. The idea is that, at the onset of learning-to-read, the status of letters changes from being independent visual objects to becoming parts of a higher-order object (i.e., the string). This is attained by adapting the mechanisms of visual object processing to the constraints of visual word processing (see Grainger, 2018; Grainger & Hannagan, 2014; see also Dehaene *et al.*, 2005, for neural correlates). Specifically, the MRF theory assumes that, with reading acquisition, location-specific letter detectors are developed and their receptive fields become reduced in size and elongated to the left – note that the initial letter is critical to translating orthographic representations into phonological representations (Grainger, Bertrand, Lété, Beyersmann, & Ziegler, 2016). Thus, the emergence of location-specific processing when learning a new script is expected to produce an initial position advantage in a TSI task.

The empirical data on the emergence of location-invariant and/or location-specific processing are very scarce (see Duñabeitia *et al.*, 2015, for an exception). Duñabeitia *et al.* (2015) examined the emergence of location-invariant processing in a longitudinal same-

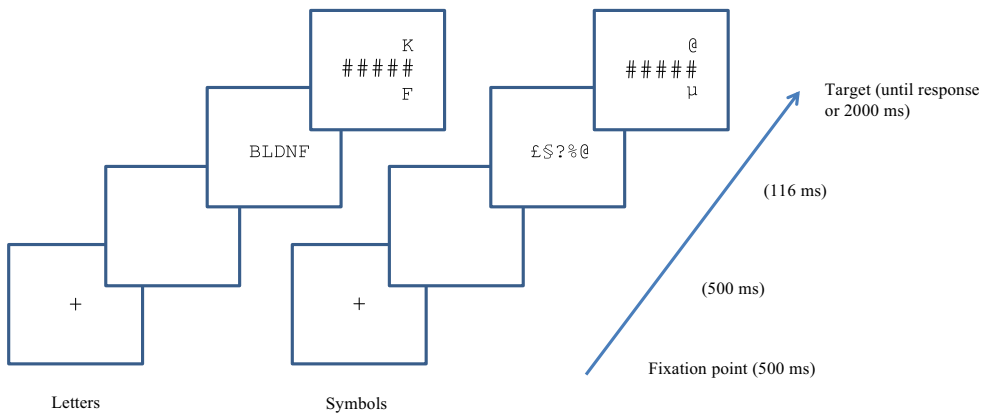


Figure 2. Depiction of the target-in-string identification task. [Colour figure can be viewed at wileyonlinelibrary.com]

different experiment with children from four years (i.e., pre-literate children) to six years (i.e., children who had acquired orthographic representations). In the accuracy data, Duñabeitia *et al.* (2015) found a significant TL effect for the older children, but not for the pre-literate children, and claimed that the ‘the skills related to the processing of internal characters’ identities and positions are inherently dependent on literacy’ (p. 548). However, *d'* (i.e., a measure of sensitivity) did not differ from zero in the experiment with pre-literate children (i.e., children were performing at chance level), which raises questions about any interpretation of the data (see Perea, Jiménez, & Gomez, 2016, for discussion)¹.

The main goal of the present experiments was to overcome this gap by examining whether acquiring a new script affects location-invariant processing and location-specific processing by using a same-different task and a TSI task, respectively. We designed a laboratory analogue of children’s reading acquisition in which adults were trained to read and write in a new, unfamiliar script. As Chetail (2017) indicated, the use of artificial scripts (i.e., sets of characters either from unknown scripts or newly devised) provides a unique opportunity to ‘examine the developmental course of a given orthographic process which is stable in adults’ (p. 103). Furthermore, recruiting adults as subjects allows us to control the participants’ prior knowledge (i.e., we can make sure that participants are not familiarized with the characters; see Maurer, Blau, Yoncheva, & McCandliss, 2010; Taylor, Davis, & Rastle, 2017), and it also enables us to increase the number of conditions and trials of the experiments: Adult participants can carry out longer experimental sessions than children, and this allows us to ensure appropriate reliability. Additionally, comparing the results of pre-literate children and developing readers may lead to intricate interpretative issues, as the accuracy and latency data vary dramatically across groups (see Perea *et al.*, 2016). Critically, the behavioural effects of learning an unfamiliar script in adults can be generalized to the effects elicited on children when learning their first language (see Taylor, Plunkett, & Nation, 2011, for discussion).

¹ While the Perea *et al.* (2016) experiment with pre-literate children rules out an interpretation of TL effects as being fully dependent on literacy acquisition, it does not provide any insights as to the emergence of location-invariant processing. To test whether location-invariant processing is influenced by literacy in young children, one would need to run a retest – ideally with letters vs. symbols (or letters from a new alphabet) – after these children learn to read.

In the current study, we employed a classic design with a pre-training phase and a post-training phase. The pre-training phase comprises two experiments: a same-different experiment on the TL effect (i.e., testing location-invariant processing; Experiment 1) and an experiment with a TSI task on the serial position function (i.e., testing location-specific processing). The pre-training experiments were conducted using eighteen consonant letters from an artificial monospaced font (BACS2serif; Vidal & Chetail, 2017). This font was used to create two different scripts: Script 1 (11 letters; two vowels and nine consonants) and Script 2 (11 letters; two vowels and nine consonants) (see Figure 3). One of the scripts was learned via print-to-sound training along five days, and the other was used as a control. An important issue is the choice of the appropriate control script. Keep in mind that the letters in the trained script would not only activate print-to-sound correspondences, but they would also be visually familiar. That is, after training, the pseudoletter ‘ϕ’ would not only correspond to a phoneme, but it also would become a familiar object. Thus, one could argue that any effects from the trained script in the post-training phase could be merely due to visual familiarity. To separate the effects of learning-to-read from the effects of visual familiarity, participants were familiarized with the visual form of the characters of the control script.

Script 1	Script 2	Phoneme
A	⊖	/a/
\	ϣ	/i/
℞	⋈	/d/
≡	┐	/k/
ϕ	└	/f/
ᐃ	⇒	/l/
ᐅ	Π	/p/
ᐆ	Λ	/m/
└	?	/r/
┐	◁	/s/
ᐇ	λ	/θ/

Figure 3. Association between the letters of the new scripts and their corresponding phonemes.

In the pre-training phase of Experiment 1, we conducted a same-different task using five-letter strings. In the pre-training phase of Experiment 2, we conducted a TSI task with five-letter strings. In both experiments, Script 1 and Script 2 were presented in separate blocks. Subsequently, each participant received training in one of the scripts (trained script) and was familiarized with the characters of the other script (control script). For the print-to-sound training, each individual learned the grapheme–phoneme associations of nine consonant letters and two vowel letters from one of the two scripts across six training sessions: Half of the participants learned the letters in Script 1 and the other half learned the letters in Script 2. Prior research has shown that readers can show some expertise in a new script quite rapidly. For instance, Chetail (2017) found that individuals acquired new regularities (e.g., letter and bigram frequency effects) after a relatively short amount of time, even in unfamiliar and complex scripts. Likewise, Brem *et al.* (2018) reported that two hours of training were enough for individuals to show some expertise for a novel script (e.g., an increase of the N1 amplitude). For the visual familiarization with the control script, participants were exposed to the eleven characters of the script, but without mentioning any orthographic or phonological information.

On day 1, participants first learned the grapheme–phoneme associations in the trained script. As in Spanish – their native language, all the grapheme–phoneme associations are transparent (e.g., the letter ‘i’ always corresponds to the phoneme/i/). Importantly, the print-to-sound training allows us to ensure that participants will learn the new script as a group of letters and not as mere symbols or shapes. As Chetail (2017) pointed out, ‘a critical feature that distinguished letters from other symbols or shapes is that letters are used to transcribe speech according to a structure code’ (p. 110). To consolidate the learning of the trained script over the next five sessions, which took place in a window of five working days, participants were asked to read aloud and write down series of items of increasing length, from 4-letter to 8-letter strings. Furthermore, the participants were familiarized with the visual forms of the characters of the control script. To that end, on each training session, each individual performed a character detection task and a character count task (see Figure 4 for depiction of each task; see also Chetail, 2017, for a similar strategy).

Finally, all participants had to pass a final test on the sixth day to show that they successfully acquired the experimental script. Once the individuals had passed this test (the criterion was set at 20 or more correct responses out of 24 in reading/writing within a time limit), they took part in the post-training experiments. What we should note here is that the print-to-sound training occurred in absence of semantics. Recent research has emphasized the role of sound-based strategies when learning-to-read a new script (e.g., see Brem *et al.*, 2018; Taylor *et al.*, 2017, for evidence with adults). This print-to-sound training also enabled us to isolate the orthographic processes from semantic processes, thus minimizing any influences from top-down processes.

The post-training experiments were parallel to the experiments from the pre-training phase. They were designed to test whether location-invariant and location-specific processing has emerged in the trained script – for comparison purposes, we also conducted a block with stimuli in an overlearned script (i.e., Roman alphabet). The predictions were clear. If location-specific orthographic processes emerge after literacy acquisition – as proposed by Dandurand *et al.* (2010) and Duñabeitia *et al.* (2014, 2015), we would expect a greater TL effect in a same-different task for the trained script in the post-training phase than in the pre-training phase. Keep in mind that, in the post-training phase, the TL effect for the newly learned script would have two constituents: a positional noise component – shared with the pre-training phase – and an orthographically based

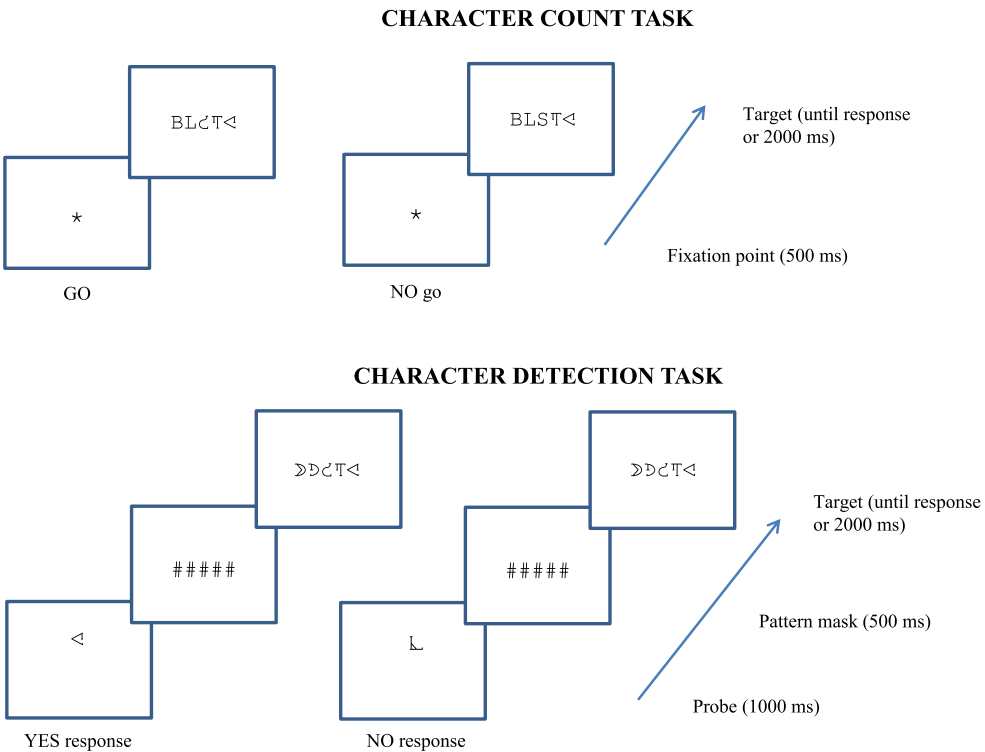


Figure 4. Depiction of the character count task (top panel) and the character detection task (low panel). [Colour figure can be viewed at wileyonlinelibrary.com]

component. For the control script, the TL effect should remain similar in magnitude in the pre- and post-training phases: the TL effect would be due to perceptual uncertainty. Alternatively, if TL effects were similar in magnitude in the pre- and post-training phase for the two scripts, this would reveal that the emergence of location-invariant processing does not emerge quickly after learning print-to-sound correspondences in a new script.

Second, if location-specific letter detectors are formed with literacy acquisition – as proposed by the MRF theory (Tydgate & Grainger, 2009), the trained script would elicit a first-letter advantage in the TSI task when measuring the serial position function in the post-training phase. This would be accompanied by an advantage of the fixed, middle position (i.e., a Λ -shape function) for the two scripts in the pre-training phase and for the control script in the post-training phase. Alternatively, if the pattern of data still shows a Λ -shape function for both the trained and untrained scripts in the post-training phase, this would suggest that print-to-sound training does not rapidly lead to the emergence of location-specific orthographic processing.

EXPERIMENT 1: LOCATION-INVARIANT PROCESSING

Method

Participants

The sample was composed of twenty-eight university students, all of them native speakers of Spanish with normal/corrected-to-normal vision and with no history of reading or

hearing disorders. All of them signed an informed consent form before participating in the experiment. Participants received a small monetary compensation after the experiment. In consonance with the registered protocol, the final number of participants was determined via a sequential Bayes factor design maximal $\underline{n}$ (see Schönbrodt & Wagenmakers, 2018, for the advantages of this approach) starting with a sample size of 28 participants. To compute the Bayes factors for the critical interaction (i.e., the three-way interaction between Phase $\times$ Script $\times$ Probe-target relationship in the accuracy data) required for the sampling procedure, we obtained the Bayes factors in the by-subjects Bayesian ANOVA – note that all the stimuli were strings of random consonants (or consonants from artificial scripts) so generalization over participants was more important than generalization over random consonants. This Bayes factor was computed in JASP (Wagenmakers *et al.*, 2018) as the ratio of the model that contained the factor of interest (i.e., all the main effects, two-way interactions, and the three-way interaction) vs. the model that did not contain the effects of interest (i.e., all the main effects and the two-way interactions). This Bayes factor exceeded 6 (i.e., the criteria established in the pre-registered protocol) in favour of the null hypothesis ($BF_{10} = .081 \rightarrow BF_{01} > 12$), so sampling was stopped with $n = 28$.

Materials

We created 240 five-consonant string pairs (probe and target) in Script 1 (see Figure 3), 240 five-consonant string pairs in Script 2 (see Figure 3), and 240 five-consonant string pairs in Roman alphabet (using Courier New font; e.g., STNGB). The two artificial scripts stemmed from the same font: BACS2serif (Vidal & Chetail, 2017). The string pairs in Script 1 and Script 2 were presented in separate, counterbalanced blocks – the string pairs in the Roman script were presented at the end of the post-training phase. All character strings were composed of non-repeated letters. There were 120 ‘different’ pairs and 120 ‘same’ pairs for each character string type. For the ‘different’ pairs in each script, 60 pairs were created by transposing two adjacent letters (e.g., $\text{CTLDV} - \text{CLTDV}$; 1st-2nd, 2nd-3rd, 3rd-4th, 4th-5th), and 60 pairs were created by replacing two adjacent letters (e.g., $\text{CTLDV} - \text{CLQDV}$; 1st-2nd, 2nd-3rd, 3rd-4th, 4th-5th). Thus, each block contained 240 pairs of character strings. In total, each participant was given 480 trials in the pre-training phase (240 string pairs in Script 1 and 240 string pairs in Script 2) and 720 trials in the post-training phase (240 string pairs in Script 1, 240 string pairs in Script 2, and 240 string pairs in Roman script). The proportion of transpositions/replacements was the same in all possible locations. To counterbalance the probe-target pairs, we created two lists for each script (see Massol *et al.*, 2013, for a similar procedure). For the practice phase, we also created eight five-consonant string pairs for each block (eight pairs in the Script 1 block, eight pairs in the Script 2 block, and eight pairs in the Roman alphabet block).

For the learning-to-read sessions, we created a template in a standard presentation software with the graphemes of the new script and the associated phoneme (see Figure 3), which were recorded by a female voice and digitalized at a sampling rate of 44.1 kHz. For each script, we created 18 items and 18 utterances of 4 characters and 5 characters, respectively, 30 items and 30 utterances of 6 characters, 66 items and 66 utterances of 7 characters, and 120 items and 120 utterances of 8 characters. The items and utterances were grouped separately by lists of 12 items that were presented with standard presentation software.

To have the participants familiarized with the characters of the control script, we employed a character detection task and a character count task (see Figure 4). For the

character detection task, we created a total of 144 pairs of items in each script (i.e., probe and target). The probe was always a single character and the target consisted of a string of artificial characters with different length (18 items of four and five characters, respectively, 30 items of 6 characters, 66 items of 7 characters and 120 items of 8 characters). For the character count task, we employed, for each script, 18 items of four and 5 characters, respectively, 30 items of 6 characters, 66 items of 7 characters and 120 items of 8 characters. Each string could be composed by artificial characters or by artificial and Roman characters.

Procedure

Pre-training-test. Participants were tested either individually or in groups of two in a quiet room. DMDX software (Forster & Forster, 2003) was used to display the sequence of stimuli and to register the timing/accuracy of the responses. Response times were measured from target onset until the participant's response. On each trial, a fixation point (*) was displayed for 500 ms in the centre of a computer screen. Next, the fixation point was replaced by a probe, which was presented for 300 ms and positioned 3 mm above the centre of the screen. Then, the target item appeared one line 3 mm below the centre of the screen. The target remained on the screen until the response or 2,000 ms had passed. All stimuli were presented in a monospaced font (15 pt BACS2serif for Scripts 1 and 2; 15 pt Courier New for the Roman letters) in black on a white background. Participants were told that they would be presented with two strings of consonants and that they would have to press the 'yes' key if they were the same, and they were asked to press the 'no' button if they were different (see Figure 1). Participants were instructed to make this decision as quickly and as accurately as possible. Eight practice trials preceded the 240 experimental trials in each block. Participants did not receive feedback during the experiment. The session lasted for 18–22 min.

Training. Participants were trained individually in the presence of the experimenter along a window of six working days in a quiet room (see Figure 5). Half of the participants learned the grapheme–phoneme associations in Script 1 and the other half learned the grapheme–phoneme associations in Script 2. There were two blocks in each session: For the trained script, participants received print-to-sound training (i.e., grapheme–phoneme association) and, for the control script, they participated in tasks that entailed visual familiarization with the control characters – the order of each block was counterbalanced across sessions.

On the first day, after the pre-training experiments, participants learned the association between the spoken forms of each grapheme in one of the unknown scripts (i.e., the experimental script: Script 1 or Script 2; see Figure 3) and they also familiarized with the visual form of the control script (i.e., the script not used for the grapheme–phoneme association). For the print-to-sound training, the characters were presented on a computer screen with their corresponding sound (participants could click on the character with the mouse and listen to its associated phoneme). Participants were also asked to hand-copy the new letters on a sheet of paper and they were given as much time as they needed to learn these associations.

For the visual familiarization part, the characters of the control script were presented one by one on a computer screen in absence of any orthographic or phonological

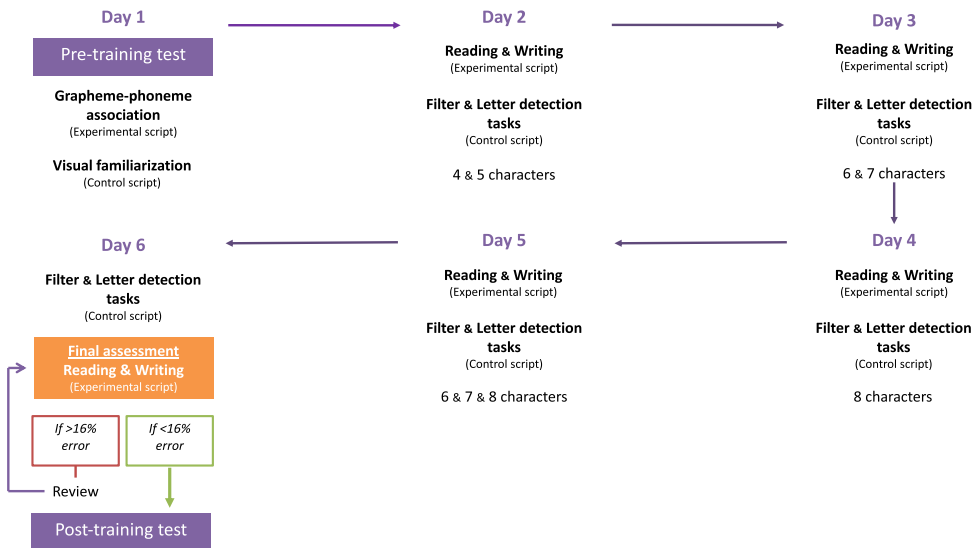


Figure 5. Schematic depiction of the training sessions. [Colour figure can be viewed at wileyonlinelibrary.com]

information. Participants had to hand-copy the control characters on a sheet of paper, without a time deadline. Then, all the characters of the control script were presented and participants took as much time as they wanted to familiarize with them (see Chetail, 2017, for a similar procedure). The experimenter checked and corrected (when necessary) the writing of the letters and characters to minimize the differences in handwriting quality².

On the second day, participants were presented with items of four and five characters; on the third day, they were presented with items of six and seven characters; on the fourth day, they practised with items of eight characters and, on the fifth day, participants were presented with items of six, seven and eight characters (see Figure 5). For the print-to-sound training, the general procedure was as follows. Participants had to read aloud and write down 36 items without time deadline. The items were presented on the computer screen, divided into alternating blocks (reading and writing) of 12 items. For reading aloud, a list of 12 items was presented on the screen (e.g., ‘ㄉㄞㄟ’) and participants were asked to read the items one by one (e.g., /daki/). During this task, the experimenter provided feedback after each item (i.e., correct/incorrect response). If the participant made a mistake, the experimenter encouraged her/him to read it again on her/his own. If the participant could not figure out the correct response, the experimenter indicated it, remarking the grapheme–phoneme correspondences. For the writing blocks, a list of 12 ‘loudspeaker’ signs (🔊) were presented on the screen. Participants were asked to press the 🔊 sign, listen to the pronunciation (e.g., /daki/), and then write down the corresponding graphemes in a sheet of paper (e.g., ‘ㄉㄞㄟ’). As blocks of 12 items were presented simultaneously, participants were able to see and listen to each item as many times as needed. The experimenter provided feedback after each item (i.e., correct/

² Exact handwritten copies of the character were not required, as neither are exact copies of the letters when children learn to write. It was enough if the handwritten copy of the character approximated to the original to be identified and distinguishable from the other characters.

incorrect response). If the participant made a mistake, the experimenter asked her/him to listen the item again. If the participant could not figure out the mistake on her/his own, the experimenter told her/him the correct response remarking the grapheme–phoneme correspondences.

For each block in both tasks (i.e., reading aloud and write down), the experimenter annotated the mistakes (if any), the order of the tasks, performing times, and other comments (e.g., the most repeated errors) in an assessment form. Moreover, after each block of 12 items, the experimenter provided general feedback of performance (i.e., correct responses, type of errors, and timing). Importantly, on the first sessions (days 2 and 3), the learning goal was to correctly establish the grapheme–phoneme correspondences. Thus, the experimenter focused mainly on the errors made by the participants. Then, on sessions 4 and 5, when the participants hardly made any mistakes, the experimenter encouraged them to read and write down as fast as possible.

For the visual familiarization block with the control script, participants completed a character detection task and a character count task in each training session. The items had the same length as the items of its corresponding print-to-sound training session. For both tasks, DMDX software (Forster & Forster, 2003) was used to display the sequence of stimuli and to register the timing/accuracy of the responses. On each trial of the character detection task (see Figure 4), a probe was presented for 1,000 ms one line 3 mm above the centre of the screen. The probe was subsequently replaced by a pattern mask with same length as the subsequent target ('#####') on the centre of the screen for 500 ms. Then, the target appeared and remained on the screen until response or 2,000 ms had passed. All stimuli were presented in a monospaced font (15 pt BACS2serif) in black on a white background. Participants were told that they would be presented with a character and then with a string of characters (both in the control script) and they would have to press the 'yes' key if the probe appeared in the subsequent string, and they were asked to press the 'no' button if the single character did not appear in the string.

On each trial of the character count task (see Figure 4), a fixation point (*) was displayed for 500 ms on the centre of a computer screen. Next, the fixation point was replaced by the target (i.e., a character string). The target remained on the screen until the response or 2,000 ms had passed. All stimuli were presented in a monospaced font (15 pt BACS2serif and 15 pt Courier New) in black on a white background. Participants were asked to press the 'yes' key only when the item presented was composed of 3 or more characters of the control script – keep in mind that the target items consisted of only control script characters (BACS2serif) or a mixture of characters of the control and Roman scripts. 30% of the items consisted of only one or two control script characters and Roman characters (i.e., participants should not press the 'yes' key). Participants received feedback on the general accuracy after each task.

Finally, on the sixth day, before conducting the post-training experiments, participants received a final test with 24 items of 8 characters (12 for reading aloud and 12 for writing) presented in the same format as in the training. They had to do the test in less than 1 min and 30 s, and 3 min and 30 s, respectively³. Those participants who passed the assessment with at least 84% of accuracy (20 out of 24 correct responses) took part in the post-training experiments⁴.

³ The time limit was set by averaging the reaction times of two pilot participants and adding 30 s more – keep in mind that the pilot participants were members of the laboratory and they were highly motivated.

⁴ A minimum of 70% accuracy was required in the visual control tasks. All participants met this criterion.

Post-training test. The post-training test was the same as in Experiment 1a, with a final additional block with Roman letters. The post-training tests lasted for approximately 25–30 min.

Results

All participants were able to write and read the newly learned script fluently and passed the final training test at the first attempt (see Appendix B, for performance of participants along the training sessions). In accordance with the pre-registered protocol, one participant was replaced because of an overall accuracy level below .60 in the replacement-letter condition.

Confirmatory analysis

The dependent variables were response time and accuracy. Error and extremely short responses (less than 250 ms: 0 responses) were omitted from the latency analyses – there were no responses longer than the 2-s deadline (i.e., they were automatically categorized as errors). The mean RTs for the correct responses and the accuracy in each experimental condition are presented in Table 1 (see also Figures 6 and 7). We performed the statistical inference not only using (generalized) linear mixed-effects models, but we also computed Bayes factors. Table 2 presents a summary of the main points of the experiment (i.e., research question, key comparisons, predictions, statistical analyses, main findings, and conclusions).

Different trials

In the inferential analyses, we focused on ‘different’ trials, as these are the ones with the TL manipulation. The main research question in the experiment was whether location-invariant processing – as measured by the TL effect – emerges in the trained but not in the untrained script in the post-training phase. To test this hypothesis, we employed (generalized) linear mixed-effects (LME) models in R (R Core Team, 2019) using the lme4 1.1-21 package (Bates, Maechler, Bolker, & Walker, 2019) and the BayesFactor 0.9.12-4.2 package (Morey & Rouder, 2018) with three fixed factors: Phase (pre- vs. post-training), Script (trained vs. untrained), and Probe-target relationship (transposed, replaced). Regarding the LME analyses, because of the normality assumption required, the raw RTs were inverse-transformed ($-1000/RT$). The most complex fitted model that converged

Table 1. Mean correct response times (in ms) and accuracy (in brackets) in the different conditions of Experiment 1

	Untrained Script			Trained Script		
	Different		Same	Different		Same
	Transposed	Replaced		Transposed	Replaced	
Pre-training	624 (.550)	602 (.714)	550 (.908)	639 (.551)	624 (.733)	581 (.868)
Post-training	585 (.603)	566 (.783)	540 (.888)	600 (.603)	572 (.795)	538 (.910)

Note. For the Roman script, the correct response times and accuracy (in brackets) were 590 ms (.529) for transposed pairs, 574 ms (.812) for replaced pairs, and 511 ms (.931) for same pairs.

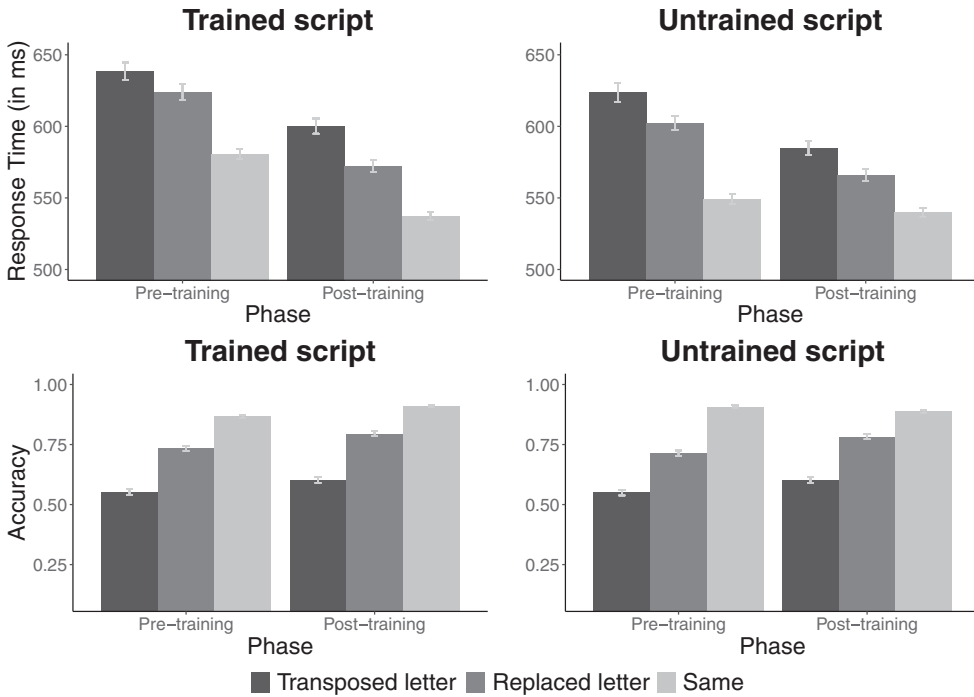


Figure 6. Mean reaction times (top panel), accuracies (bottom panel), and standard errors in the trained and untrained scripts of Experiment I.

was: $-1000/RT \sim \text{script} * \text{condition} * \text{phase} + (1 + \text{script} + \text{phase} | \text{subject}) + (1 | \text{item})$. For the generalized linear mixed analyses of the accuracy data, responses were coded as binary values (1 = correct, 0 = incorrect) and we used the `glmer` function in the `lme4` package (family = binomial). The most complex fitted model that converged was as follows: $\text{accuracy} \sim \text{script} * \text{condition} * \text{phase} + (1 + \text{script} | \text{subject}) + (1 | \text{item})$. (In Appendix A, we report the [non-pre-registered] analyses using Bayesian linear mixed-effects models with the maximal random structure – the results were essentially the same as those reported here.) To compute the Bayes factors on the latency data, we used `lmBF` function from the `BayesFactor` package with the default Cauchy distribution (centred around 0 and with a width parameter $\delta = 0.707$) (see Rouder, Speckman, Sun, Morey, & Iverson, 2009; Wagenmakers *et al.*, 2017, for discussion). To compute the Bayes factors on the accuracy data, we calculated the Bayes factors from Bayesian analyses of variance (ANOVAs) with the aggregated data by participants – note that items were strings of letters in an artificial script. For the computation of the Bayes factors for each effect, we followed the same logic as in prior research (see Leininger, Myslín, Rayner, & Levy, 2017; Staub & Goddard, 2019, for illustration). For the numerator, we compared the maximal model that included the effect of interest vs. a null model that does not assume any fixed effects or interactions. For the denominator, we compared the maximal model after excluding the effect of interest vs. a null model that does not assume any fixed effects or interactions. The ratio between these two Bayes factors was the Bayes Factor of the effect.

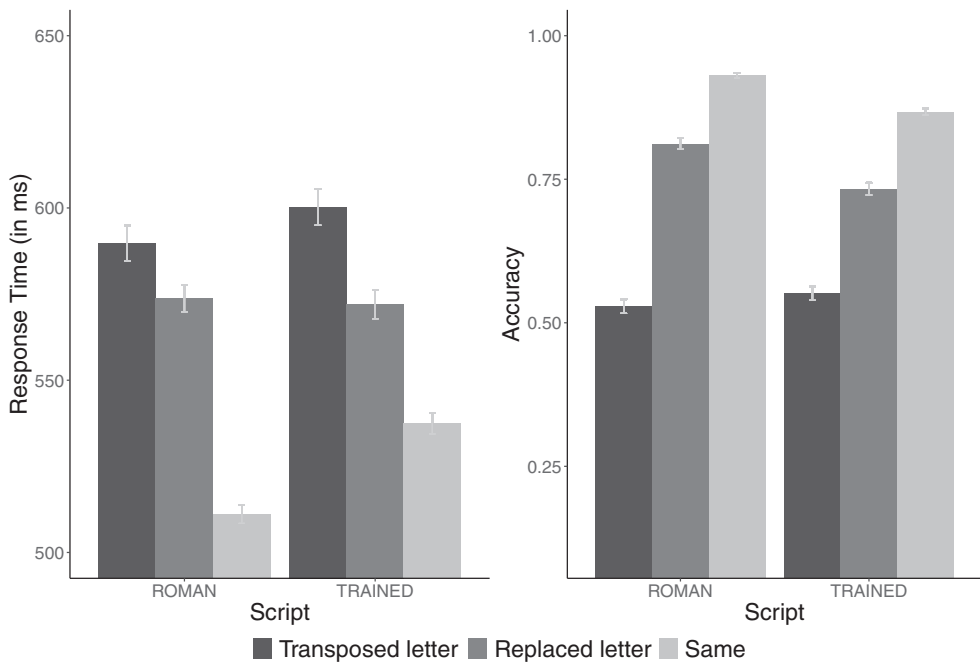


Figure 7. Mean reaction times (left panel), accuracies (right panel), and standard errors in the trained (post-training phase) and Roman scripts of Experiment 1.

Latency analyses. Responses were, on average, 21 ms faster for the RL condition than for the TL condition (i.e., overall TL effect; 591 vs. 612 ms; $b = .08$; $t = 6.18$ $p < .001$; $BF_{10} = 6.69e + 05$) and participants responded, on average, 41 ms faster in the post-training tests than in the pre-training tests (581 vs. 622 ms; $b = .13$; $t = 4.55$ $p < .001$; $BF_{10} = 20.34$). We found no overall differences in response times between the trained and untrained script ($b = -.01$; $t = -.60$, $p = .55$; $BF_{10} = 1.48$). The interaction between Phase and Script barely reached the significance level in the frequentist analyses ($b = -.04$; $t = -2.43$, $p = .02$), but the Bayes factors indicated anecdotal evidence towards a null effect ($BF_{10} = 0.55$). None of the other interactions approached significance (all $ts < .78$, $ps > .59$; all $BF_{10} < .35$) (see Figure 6).

Accuracy analyses. Accuracy was higher for the RL condition than for the TL condition (.756 vs. .577; $b = -1.05$; $z = -12.73$, $p < .001$; $BF_{10} = 1.071e + 39$), and participants were more accurate in the post-training phase than in the pre-training phase (.696 vs. .637; $b = -.38$; $z = -4.44$, $p < .001$; $BF_{10} = 415911$). Neither the effect of Script ($b = -.08$; $z = -.83$, $p = .41$; $BF_{10} = .21$) nor any of the interactions approached significance (all $zs < 1.19$, $ps > .23$; all $BF_{10} < .26$) (see Figure 6).

Exploratory analyses

As indicated in the pre-registration, we also compared the TL effect of the trained script (post-training phase) and the TL effect in an overlearned script (i.e., the Roman script). The two fixed factors in the analyses were Script (Trained [post-training] vs. Roman) and Probe-target relationship (transposed, replaced) – the inferential analyses were parallel to

Table 2. Summary of Experiment I (Location-invariant processing: same-different task)

Question	Key comparisons	Predictions	Confirmatory (Different trials)	Statistical analyses	Results (main findings)	Conclusions
Does location - invariant processing emerge rapidly after learning-to-read an artificial script?	Script (Trained vs. Untrained) ×	Greater TL effect for the trained script in the post-training phase than in the pre-training phase.	✓	(G)LME — 1000/RT ~ script*condition*phase + (1 + script+phase subject) + (1 item)	Condition Phase	The magnitude of the transposed- letter effect was similar for the trained script and for the untrained script. Location-invariant processing does not emerge rapidly after learning-to-read an artificial script
	Phase (Pre- vs. Post-training) ×		✓	BF accuracy ~ script*condition*phase + (1 + script subject) + (1 item) RT ~ subject * script * condition * phase	Condition*Script Phase Condition Phase	
	Condition (Transposed vs. Replaced letter)	Similar TL effect for the control script in both phases.	×	BRMS Accuracy: Bayesian ANOVA with script, condition, and phase as factors 1000/RT ~ script * condition * phase + (1 + script * condition * phase subject) + (1 + condition * phase item)	Condition Phase Condition Phase	
				accuracy ~ script * condition * phase + (1 + script * condition * phase subject) + (1 + condition * phase item)	Condition Phase	
	Script (Trained [post-training] vs. Roman) ×	Greater TL effect for the Roman script than for the trained script.	Exploratory	(G)LME — 1000/RT ~ script*condition + (1 + script subject) + (1 item) accuracy ~ script*condition + (1 + script subject) + (1 item)	Condition Condition	The magnitude of the transposed- letter effect was greater in the Roman script than in the newly learned script in the accuracy analyses
	Condition (Transposed vs. Replaced letter)		✓	BF RT ~ subject * script * condition * phase Accuracy: Bayesian ANOVA with script and condition as factors	Condition*Script Condition Condition*Script Condition	
			×	BRMS — 1000/RT ~ script*condition + (1 + script*condition subject) + (1 + script*condition item) — 1000/RT ~ script*condition + (1 + script*condition subject) + (1 + script*condition item)	Condition Condition*Script	
					Condition Condition*Script	

Continued

Table 2. (Continued)

Question	Key comparisons	Predictions	Same trials	χ^2	LME	Statistical analyses	Results (main findings)	Conclusions
Script (Trained vs. Untrained) $\times$ Phase (Pre- vs. Post-training)			Same trials	χ^2	LME	-1000/RT ~ script*phase + (1 + phase subject) + (1 item) accuracy ~ script*phase + (1 + phase subject) + (1 item)	Phase Script*Phase Phase Script Script*Phase Phase Script*Phase Phase Script Script*Phase	Learning-to-read in the new script helped encoding the letter strings emergence of rudimentary orthographic representations. Same' responses were substantially faster and more accurate in the post-training phase than in the pre-training phase only for the trained script.

Note. ✓ [pre-registered analyses]. χ^2 [non-pre-registered analyses]; number of participants was determined via sequential Bayes factor design obtaining BF in the by-subjects ANOVA ($BF_{10} = .081 \rightarrow BF_{01} > 12$).

those described above. The most complex fitted model that converged was: *Dependent_Variable* [$-1000/\text{RT}$ or accuracy] $\sim \text{script} * \text{condition} + (1 + \text{script} | \text{subject}) + (1 | \text{item})$ (see Figure 7).

Latency analyses. Participants responded, on average, 22 ms faster in the RL condition than in the TL condition (573 vs. 595; $b = .077$ $t = 6.00$, $p < .001$; $\text{BF}_{10} = 129.80$). There were no overall differences between the trained script and the Roman script ($b = .02$ $t = .38$, $p = .71$; $\text{BF}_{10} = 0.53$). The interaction between probe-target relationship and script was not significant ($b = -.01$; $t = -.41$, $p = .68$; $\text{BF}_{10} = 0.54$).

Accuracy analyses. Participants were more accurate in the RL condition than in the TL condition (.804 vs. .566; $b = -1.12$; $z = -13.06$, $p < .001$; $\text{BF}_{10} = 3.308e + 22$), whereas the effect of script was not significant ($b = .22$; $z = 1.14$, $p = .25$; $\text{BF}_{10} = 0.734$). We found a significant interaction between the two factors ($b = -.58$; $z = -4.67$, $p < .001$; $\text{BF}_{10} = 7.758$), which reflects that the TL effect was greater in the Roman script than in the trained script (.283 vs. .192).

Same trials

While not indicated in the pre-registered protocol, the examination of ‘same’ responses in the pre- and post-test phases for the trained and untrained scripts may shed some light on the role of orthographic-phonological training when processing letter strings. To analyse ‘same’ responses, we employed (generalized) linear mixed-effects models on the latency and accuracy data. The two fixed effects were Script (trained vs. untrained) and Phase (pre-training, post-training). The most complex model that converged in the (generalized) linear mixed-effects models was: *Dependent_Variable* [$-1000/\text{RT}$ or accuracy] $\sim \text{script} * \text{phase} + (1 + \text{phase} | \text{subject}) + (1 | \text{item})$. These analyses were complemented with Bayesian linear mixed-effects models using the maximal random factor structure (see Appendix A).

Latency analyses. Responses were, on average, 26 ms faster in the post-training phase than in the pre-training phase (539 vs. 565 ms; $b = .14$; $t = 2.93$, $p = .01$), whereas there were no signs of an effect of script ($b = .01$; $t = -.51$, $p = .61$). More important, the interaction between Script and Phase was significant ($b = -.12$; $t = -8.09$, $p < .001$). This reflected that responses were faster in the post-training than in the pre-training phase for the trained script (41 ms; 537 vs. 581 ms), but not for the untrained script (9 ms; 540 vs. 549 ms) (see Figure 6).

Accuracy analyses. Participants were more accurate in the post-training than in the pre-training phase (i.e., main effect of phase; $b = -.48$; $z = -2.97$, $p = .003$) and with the untrained than with the trained script (i.e., main effect of script; $b = -.27$; $z = -3.17$, $p = .001$) – note that the effect of script was .009 and was not corroborated by the Bayesian linear mixed-effects analyses (see Table A3). More important, mimicking the latency analyses, we found an interaction between the two factors ($b = .72$; $z = 6.16$, $p < .001$): Participants were more accurate in the post-training test than in the pre-

training test for the trained script (.910 vs. .868, respectively), but not for the untrained script (.888 in the post-training vs. .908 in the pre-training test) (see Figure 6).

Discussion

As usual, we found a substantial transposed-letter effect for ‘different’ responses in the same-different task: Participants’ responses were faster and more accurate for replacement-letter pairs than for transposed-letter pairs (see Krueger, 1978; Ratcliff, 1981, for early evidence; see also Duñabeitia *et al.*, 2012; Massol *et al.*, 2013; Perea *et al.*, 2016). More important for the purposes of the experiment, the magnitude of the transposed-letter effect was similar for the trained script and for the (visual) control script in both latency and accuracy data. We did find that the responses to ‘different’ trials were, on average, faster and more accurate in the post-training phase than in the pre-training phase. However, this occurred similarly in both the trained and visual control scripts; hence, it could have been due to the participants’ being more visually familiar with the new letters. In addition, the size of the transposed-letter effect was greater in the Roman script than in the newly learned script in the accuracy analyses (28.3% vs. 19.2% of errors, respectively) – note that previous studies on the transposed-letter effect also showed significant effects on accuracy, but not on response latencies (e.g., Massol *et al.*, 2013; Perea *et al.*, 2016; Perea & Lupker, 2004)⁵. This is consistent with the idea of orthographic location-invariant mechanisms being at work in the Roman script, but not in the newly learned script (see Duñabeitia *et al.*, 2012; Massol, *et al.*, 2013; see also García-Orza *et al.*, 2010; Muñoz *et al.*, 2012, for greater transposed-letter effect for letters than for other visual objects [symbols, digits, false fonts]).

The above results may offer the impression that training a new script did not create any stable orthographic representations. However, this interpretation is difficult to reconcile with the fact that ‘same’ responses were substantially faster and more accurate in the post-training phase than in the pre-training phase for the trained script (538 vs. 581 ms; .910 vs. .868), but not for the untrained script (540 vs. 549 ms; .888 vs. .908). This finding strongly suggests that learning-to-read in the new script helped encoding the letter strings, thus reflecting the emergence of rudimentary orthographic representations.

In sum, the current same-different experiment favours the view that location-invariant processing, as measured by the transposed-letter effect, does not emerge rapidly after learning-to-read in a new script. We defer a more detailed discussion of this issue in the General Discussion.

EXPERIMENT 2: LOCATION-SPECIFIC PROCESSING

Method

Participants

They were the same as in Experiment 1. To compute the Bayes factors for the critical interaction (i.e., the three-way interaction between Phase $\times$ Script $\times$ Position in the

⁵ Furthermore, for the Roman script, we found that size of the transposed-letter effect was considerably smaller for external than for internal transpositions: 14.5% vs. 42.1%, respectively (e.g., see Gomez *et al.*, 2008, for a similar pattern). In contrast, for the trained script, the size of the transposed-letter effect was only slightly lower for external transpositions than for internal transposition (17.0% vs. 21.4%, respectively). This again suggests that the transposed-letter effect in the Roman script and the trained script reflects different underlying processes.

accuracy data) required for the sampling procedure, we obtained the Bayes factors in the by-subjects Bayesian ANOVA. This Bayes factor exceeded 6 (i.e., the criteria established in the pre-registered protocol), $BF_{10} = .05 \rightarrow BF_{01} = 20$, so sampling was stopped with $n = 28$.

Materials

Based on the design and procedure used by Tydgate and Grainger (2009), we created a set of 180 five-consonant strings in two unfamiliar scripts: 90 in Script 1 and 90 in Script 2, and – for the post-training phase – a set of 90 five-consonant strings in Roman alphabet (e.g., STNGB). None of the letter strings contained repeated characters.

We designed three different blocks (one for script: Script 1, Script 2, and Roman script) with 90 experimental and 9 practice trials each one. The order of the artificial script blocks was counterbalanced between subjects. We manipulated the target position in the array (1st, 2nd, 3rd, 4th, and 5th position). As in the Tydgate and Grainger (2009) experiments, each of the target characters was presented 2 times at each of the five target positions (once above and once below the backward mask), and 40 times at a non-target position (i.e., each target character played as alternative at each of the five positions). Importantly, the incorrect alternative was never presented in the stimulus array. We created two lists for each of the artificial scripts, manipulating the orientation of the target character (i.e., in List 1, the target was presented above the array, whereas in List 2 the same target was presented below the array). These two lists were presented to all participants, one for the pre-training test and the other for the post-training test. The sub-experiment with the Roman script was presented at the end of the post-training test.

The learning sessions materials were the same as in Experiment 1.

Procedure

Pre-training test. Participants were tested individually or in groups of two in a quiet room. DMDX software (Forster & Forster, 2003) was used to display the sequence of stimuli and to record the timing and accuracy of the responses. Each trial began with a fixation point (i.e., '+') that stayed on the screen for 500 ms and was followed by a 500-ms with the blank screen. Then, a string of five letters was presented for 116 ms (see Scaltritti *et al.*, 2018, for the same set-up). The array of characters was followed by a backward mask ('#####') accompanied by two characters, above and below the mask, respectively, at one of the five possible array positions (i.e., characters position as a post-cue) (see Figure 2). The stimuli were displayed on the screen until the participant responded or 2 seconds had passed. All stimuli were presented in black on a white background. We employed a monospaced font for the two scripts. Participants were asked to decide which of the two characters was present in the corresponding position of the preceding array. They were required to press the 'up arrow' key on the keyboard for the character above and the 'down arrow' key for the character below the array. They were explicitly instructed to fixate at the centre of the array and make the decision as quickly and as accurately as possible. The two scripts were presented in separated blocks, counterbalanced by subjects. Nine practice trials preceded the 90 experimental trials in each of the experimental conditions (90 trials in Script 1 and 90 trials in Script 2). Participants did not

receive feedback during the experiment. There was a short break between blocks. The session lasted for around 20–25 min.

Training. It was the same as in Experiment 1 (see Figure 5).

Post-training-test. It was the same as in the pre-training test, except for the addition of a third block with stimuli in the Roman script. The session lasted for around 30 min.

Results

As indicated in the pre-registration, those participants with less than .60 of accuracy in the middle position in the pre- and post-training experiments were replaced – this occurred with four participants. Mean accuracies (and standard errors) for all target types and target positions are presented in Figure 8. The three fixed factors were Phase (pre- vs. post-training), Script (trained vs. control), and Position (1st, 2nd, 3rd, 4th, 5th). By-subjects and by-items classical and Bayesian ANOVAs were performed on the accuracy data. The computation of the Bayes factors was parallel to that described for accuracy in Experiment 1. In Appendix A, we report supplementary [non-pre-registered] analyses using Bayesian linear mixed-effects models (see Table 3 for a summary of the main points of Experiment 2).

Confirmatory analyses

The ANOVAs showed that accuracy was a function of serial position, $F_1(4, 108) = 67.87$, $p < .001$, $BF_{10} = 2.497e + 66$; $F_2(4, 175) = 106.45$, $p < .001$, $BF_{10} = 9.445e + 36$. Accuracy levels were higher in the central, third position (.778) than in the first, second, fourth, and fifth positions (with mean accuracy levels of .535, .502, .510, and .507, respectively). Furthermore, the overall accuracy levels were virtually the same for the pre- and post-training phase, $F_1(1, 27) = 1.18$, $p = .29$, $BF_{10} = .10$; $F_2 < .001$, $BF_{10} = .09$, and for the trained and control scripts (both F s < 1 ; both BF s $< .11$) (see Figure 6).

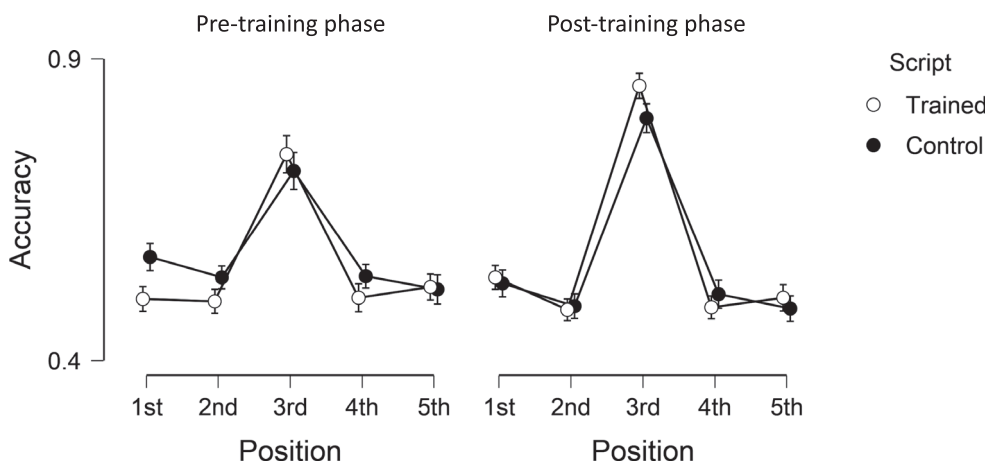


Figure 8. Mean accuracies and standard errors in the trained and untrained scripts of Experiment 2.

Table 3. Summary of Experiment 2 (location-specific processing; target-in-string identification task)

Question	Key comparisons	Predictions	Statistical analyses	Results (main effects)	Conclusions	Location-specific
Does location-specific processing emerge rapidly after learning-to-read an artificial script?	Script (Trained vs. Untrained)	Greater accuracy for first-letter position in the trained script than in the post-training phase	Confirmatory	Position (3 rd)	Advantage of the middle position	Advantage of the middle position over the other positions for the trained and untrained scripts in the pre- and post-training phases
	Phase (Pre- vs. Post-training)	phase – this would be accompanied by the advantage of the middle position.		Position (3 rd)*Phase	Accuracy on third position greater in the post-training phase than in the pre-training phase	
	×			Position (3 rd)		
	Position (1 st , 2 nd , 3 rd , 4 th , 5 th)			Position (3 rd)*Phase		
	Script (Trained vs. [post-training] vs. Roman)	Greater first-letter advantage for the Roman script than for the trained script (post-training phase)	Exploratory	Position (3 rd)	Advantage of the middle position	Attentional capture to the centre of the string
	×			Position (3 rd)	No differences between the Roman and Trained script	
	Position (1 st , 2 nd , 3 rd , 4 th , 5 th)			Position (3 rd)		
				Position (3 rd)*Phase		

Note. ✓ [pre-registered analyses]. ✗ [non-pre-registered analyses]; number of participants was determined via sequential Bayes factor design obtaining BF in the by-subjects ANOVA (BF₁₀ = .05 → BF₀₁ = 20)

The serial position function differed in the pre-training and post-training phase (Position $\times$ Phase interaction; $F_1(4, 108) = 8.65$, $p < .001$, $BF_{10} = 241.21$; $F_2(4, 175) = 6.22$, $p < .001$, $BF_{10} = 53.42$): This occurred because accuracy in the third position was higher in the post-training phase than in the pre-training phase (.828 vs .728, respectively). Neither the interaction between Position and Script ($F_1(4, 108) = 2.65$, $p = .04$, $BF_{10} = .23$; $F_2 < 1$, $BF_{10} = .01$) nor the interaction between Script and Phase ($F_1(1, 108) = 3.22$, $p = .08$, $BF_{10} = .50$; $F_2 < 1$, $BF_{10} = .12$) were significant. Finally, there were no signs of a Phase $\times$ Script $\times$ Position interaction (both $F_s < 1$; $BF_{10} < .09$).

Exploratory analyses

We also compared the serial position function of the trained script (in the post-training phase) and an overlearned script (i.e., Roman script) (see Figure 9). The two fixed factors in the ANOVAs were Script (Trained [post-training] vs. Roman) and Position (1st, 2nd, 3rd, 4th, 5th). The analyses were parallel to those described above.

Accuracy was a function of serial position, $F_1(4, 108) = 154.29$, $p < .001$, $BF_{10} = 4.104e + 48$; $F_2(4, 260) = 60.82$, $p < .001$, $BF_{10} = 7.918e + 23$. Participants were substantially more accurate on the third position (.841) than on the other letter positions (.540, .504, .525, and .527, in the first, second, fourth, and fifth positions, respectively). In addition, we did not find any clear signs of a difference in the overall accuracy levels in the trained script and the Roman script (.574 vs. .601; $F_1(1, 27) = 3.14$, $p = .09$; $BF_{10} = .17$; $F_2 < 1$; $BF_{10} = .78$). Finally, as can be seen in Figure 9, the serial position functions of the Roman and trained scripts were remarkably similar and the interaction between the two factors was not significant, $F_1(4, 108) = 2.28$, $p = .09$; $BF_{10} = .09$; $F_2(4, 260) = 2.05$, $p = .09$, $BF_{10} = .31$).

Discussion

The current experiment, using a target-in-string identification task, showed an advantage of the middle, fixated position over the other positions for the trained and control scripts

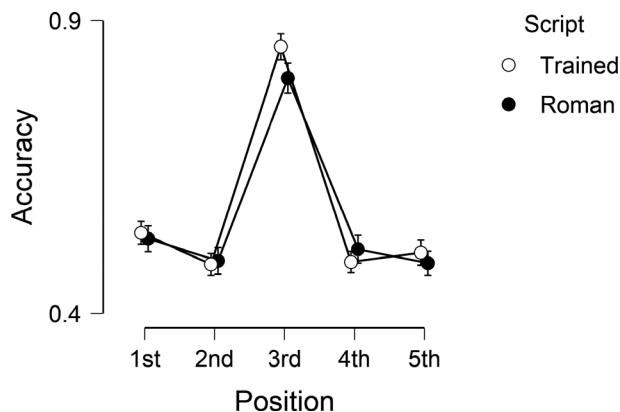


Figure 9. Mean accuracies and standard errors in the trained (post-training phase) and Roman scripts of Experiment 2.

not only in the pre-training phase, but also in the post-training phase. We did find a small numerical advantage of the initial position over the second position (see Figure 8), but this difference was similar for the trained and untrained script in the pre-/post-training phases – this pattern was also corroborated in the analyses with Bayesian linear mixed-effects models (see Appendix A). We also found that accuracy was higher in the post-training than in the pre-training session – this occurred mainly in the central, fixated letter. This effect was similar for the trained script and for the visual control script; thereby, it can be parsimoniously explained in terms of visual familiarity. Taken together, these findings strongly suggest that location-specific letter detectors do not emerge rapidly after learning-to-read and write in an artificial script.

Finally, in the exploratory analyses, we failed to find a stronger initial-letter advantage in the Roman script when compared to the trained script. We prefer to keep cautious about this latter finding. First, in the instructions, we emphasized that participants should be looking at the fixation point at the beginning of each trial. Second, because it was an exploratory analysis, participants always performed the task with Roman letters in the last block. As a result, the attentional capture at the middle position that occurred in initial blocks with the artificial scripts could have been dragged into the last block. A more comprehensive explanation is presented in the General Discussion section below.

GENERAL DISCUSSION

We designed two experiments to track the emergence of early orthographic processes in the first stages of learning-to-read through the examination of two markers of orthographic processing (Grainger, 2018): location-invariant processing (Experiment 1) and location-specific processing (Experiment 2). To that aim, we employed a design with a pre-training phase and a post-training phase in which adults were trained in reading and writing nonsense words in an artificial script along six sessions. Participants successfully mastered the trained script after the learning-to-read training (see Appendix B). All of them passed the final assessment in the prescribed time with no errors in the reading aloud and writing down tasks. To control for visual familiarity, participants were also familiarized with the visual form of the characters of the control script during the training sessions.

The emergence of location-invariant processing

The first research question was whether readers show some location-invariant processing for the newly learned script on top of the position uncertainty that may affect all visual objects in a string. As stated in the Introduction, location-invariant processing has been proposed to emerge with literacy acquisition (Dandurand *et al.*, 2010; Duñabeitia *et al.*, 2014, 2015).

To our knowledge, the only published study that directly examined this issue was conducted by Duñabeitia *et al.* (2015). They employed a longitudinal design using a same-different experiment in which a group of children was tested in their antepenultimate pre-school year (i.e., pre-literate children; mean age = 4.24 years), in their last pre-school year (i.e., pre-literate children; mean age = 5.21 years), and in the first year of primary school (i.e., they had already acquired literacy skills; mean age = 6.32 years). They used four-letter strings in which, for ‘different’ trials, they had transposed-letter and replaced-letter pairs. Duñabeitia *et al.* (2015) only found a transposed-letter effect when the children were in first grade (42.9% vs. 30.6% of errors, for transposed vs. replaced-letter pairs) and

concluded that ‘position uncertainty emerges as a consequence of literacy training’ (p. 549). However, the null effects obtained with the children in their pre-school years faced interpretative difficulties because these children did not adequately perform the task (d' was close to zero; see Perea *et al.*, 2016). Indeed, Perea *et al.* (2016) found robust transposed-letter effects in pre-literate children with a simplified version of the same-different task. Nevertheless, they did not run a retest when these children learned to read – note that first graders could have been at ceiling with this simplified task, thus making the comparison uninterpretable.

To test the emergence of location-invariant processing, while avoiding the interpretive difficulties of comparing performance of pre-school vs. school children, we conducted a same-different matching task with adult participants using transposed-letter vs. replaced-letter pairs as ‘different’ trials before and after learning-to-read in a new script (Experiment 1). To control for mere visual familiarity, we included an untrained script that was also presented during training. Results showed transposed-letter effects of similar size for the trained and untrained scripts in both the pre- and post-training phases. Furthermore, Bayesian analyses offered substantial evidence in favour of a null interaction between training, phase, and script. Thus, learning-to-read and write in a new script does not lead to the rapid emergence of location-invariant processing.

Importantly, we did find some training benefit in the trained script for ‘same’ responses in both response times and accuracy, which suggests the emergence of an early and basic visual specialization for letter strings. However, the newly acquired expertise in the new script was not sufficient to induce location-invariant processing. Indeed, the transposed-letter effect was greater for the Roman script than for the trained script (i.e., 28.3% vs. 19.2% of errors, respectively). This is the typical pattern when comparing strings of letters vs. strings of other visual objects (e.g., symbols, unknown letters)⁶. This pattern can be parsimoniously explained in terms of an orthographically specific location-invariant component in the Roman script over and above the location uncertainty common to all visual objects (see Massol *et al.*, 2013).

Our findings can also shed some light on the early developmental trajectory of the letter-specific position coding when learning-to-read. The absence of the emergence of location-invariant processing in the very early stages of learning-to-read can be accommodated by the dual-route model of orthographic development proposed by Grainger and Ziegler (2011; see also Grainger *et al.*, 2012; Ziegler, Bertrand, Lété, & Grainger, 2014). This model assumes that, in the first stages of reading acquisition, the processing of letters in a word is serial, thereby letter position coding is very strict (i.e., fine-grained orthographic coding). It is only when readers have more extensive reading experience that a more parallel processing of letters is developed, thus speeding the mapping of letters onto orthographic representations and producing greater transposed-letter effects (i.e., coarse-grained orthographic coding; see Grainger *et al.*, 2012). In the context of the current experiment, participants acquired some basic orthographic skills, as revealed by better performance for ‘same’ responses in the post-training phase. However, this expertise did not suffice for a coarse-grained processing to emerge. Indeed, in a lexical decision experiment that compared the error rates to pseudohomophone and

⁶ Furthermore, for the Roman script, we found that size of the transposed-letter effect was considerably smaller for external than for internal transpositions: 14.5% vs. 42.1%, respectively (e.g., see Gomez *et al.*, 2008, for a similar pattern). In contrast, for the trained script, the size of the transposed-letter effect was only slightly lower for external transpositions than for internal transposition (17.0% vs. 21.4%, respectively). This again suggests that the transposed-letter effect in the Roman script and the trained script reflects different underlying processes.

orthographic controls, Grainger *et al.* (2012) found effects greater than 30% in Grade 1 and Grade 2 children – these effects were smaller with older children (i.e., the effects were 20% in Grade 3, 21% in Grade 4 and 16% in Grade 5). That is, beginning readers use phonological recoding (i.e., a fine-grained orthographic coding) rather than the coarse-grained coding responsible for location-invariant processing. Thus, the greater transposed-letter effect in the Roman (overlearned) script than in the newly learned script obtained in the current experiments suggests that the emergence of location-invariant processing requires a more complete establishment of a written orthographic code (see Grainger *et al.*, 2012; Grainger & Ziegler, 2011; Ziegler *et al.*, 2014)⁷.

The emergence of location-specific processing

The second research question was whether location-specific processing emerges rapidly for the newly learned script using a target-in-string identification task. The dissociation in the accuracy serial position functions of letters (W-shape function) and symbols (Λ-shape function) in this task is assumed to be to an adaptation of the mechanisms of visual object processing to cope with visual word processing (i.e., location-specific letter detectors creation). Importantly, Dandurand *et al.* (2010; Grainger *et al.*, 2016) hypothesized that the conversion of the mechanisms of simple visual object processing into location-specific detectors occurs with reading acquisition.

Results in the target-in-string identification task (Experiment 2) showed a clear advantage of the middle position for both the trained and control scripts in all scenarios. More critically, we found no signs of an interaction in accuracy between training, phase, and position, as shown in the Bayesian analyses. We also found a small advantage of the initial-letter position over the other letter positions (see Figure 6; see also Appendix A), but this difference was not modulated by training (i.e., the difference was approximately constant in the pre- and post-training phases and in the trained and untrained scripts). Finally, the overall accuracy in the post-training phase was greater than in the pre-training phase for both, the learned and the control script, but this occurred essentially for the middle, fixated, letter. This latter finding can be parsimoniously explained in terms of better performance due to increased visual familiarity rather than on location-specific processing.

To our knowledge, unlike for location-invariant processing, no study has directly examined the emergence of location-specific processing – neither with children nor with adults. Nevertheless, for comparison purposes, it may be relevant to briefly discuss the studies that examined the developmental trajectory of the first-letter advantage in the very early stages of learning-to-read. Grainger *et al.* (2016) showed a small increase in the initial-letter advantage across school grades (from 1st to 4th) (see also Schubert, Badcock, & Kohnen, 2017, for a similar pattern of results). Importantly, the accuracy in the first and second position for 1st- to 3rd-grade children was very similar, around 55% and 60% in the Grainger *et al.*'s (2016) experiment. The lack of a sizeable first position advantage in the initial grades in developing readers is in consonance with the results of Experiment 2. Taken together, these findings suggest that location-specific processing does not emerge

⁷ An alternative account of orthographic development is Castles *et al.*'s (2007) lexical tuning model. The model assumes that acquiring more and more words in the lexicon involves an increasingly dense neighbourhood of orthographically similar words. To efficiently identify these words, the orthographic representations become increasingly fine-tuned – this includes more precise positional representation of the visual input. Our experiments, however, were not designed to test the development of the orthographic lexicon (i.e., participants were trained to read and write pseudowords).

in the first stages of learning-to-read; instead, there appears to be a long route for this mechanism to emerge and develop.

Nevertheless, the lack of the sizeable first-letter advantage in the (overlearned) Roman script suggests that some caution when interpreting the findings of Experiment 2. In the experimental set-up, participants always received the blocks with the artificial scripts (i.e., the main blocks for the purposes of the experiment) before the final block with the Roman script, and furthermore, instructions stressed that they should fixate at the centre position at the beginning of each trial. Thus, a parsimonious explanation of the lack of a substantial first-letter advantage in the Roman script is that, to cope with the highly demanding blocks with the artificial scripts, participants' attention was focused on the middle position and this strategy was dragged into the Roman block. Thus, one might argue that the settings of Roman block were not optimal to capture a W-serial position function in the Roman block. Future research should examine to what degree the accuracy function in target-in-string identification tasks is modulated by task context and instructions (e.g., see Winskel *et al.*, 2014, for evidence of different accuracy functions depending on the nature of the writing system).

On the emergence of orthographic processing when learning-to-read a new script

Recent research with adult readers has shown that orthographic processes can emerge rapidly after learning a script during a relatively short amount of time (e.g., Chetail, 2017; Lally, Taylor, Lee & Rastle, 2020; Taylor *et al.*, 2011). For instance, in the Taylor *et al.*'s (2011) experiments, adults learned 36 words in an artificial script during 30–45 min. In the post-training phase, participants had to discriminate between trained and untrained items (i.e., an analogue to lexical decision). Results showed that participants could successfully discriminate trained from untrained items and, more important, the response times to the trained items were sensitive to vowel frequency. In a subsequent generalization phase, participants were asked to read aloud a series of new (untrained) items. Results showed an effect of both vowel frequency and consistency. All and all, the Taylor *et al.* (2011) experiments suggest that participants can quickly and efficiently extract sub-word spelling–sound regularities in a new script (see Chetail, 2017, for a similar pattern of results regarding letter and bigram frequency).

More recently, Lally *et al.* (2020) conducted an experiment in which participants learned 24 five-letter pseudowords either in a sparse or in a dense artificial orthography, using a between-subject design, during a four-day training. (The 24 pseudowords in the dense orthography included 12 anagram pairs, whereas none of the 24 pseudowords sparse included anagrams.) When tested in an old–new recognition task (i.e., they had to discriminate between trained and untrained items), participants made fewer false positives for untrained items created by transposing two letters in the dense orthography than in the sparse orthography. Therefore, the findings reported by Lally *et al.* (2020) are a demonstration that the properties of the writing systems may modulate how letter order is encoded in a newly learned script.

In the current experiments, participants were able to read and write in the trained script with some fluency. Notably, in line with above-cited studies with artificial script training, we found some letter-specific processing as a consequence of learning-to-read: responses to 'same' trials in the same-different task were faster and more accurate in the post-trained phase for the trained script, but not for the untrained script. As Krueger (1978) indicated, fast and accurate responses for 'same' trials imply that participants require an exhaustive processing of the letter string. Thus, this pattern suggests that early

literacy induced some specialization for letter strings that made the identification of same pairs more effortless. However, we found no signs reflecting the emergence of location-invariant and location-specific processing in the newly learned script.

In addition, we found an improvement in performance in the post-training phase for both the learned script and the visual control script that can be parsimoniously explained in terms of visual familiarity. Indeed, previous event-related potentials (ERP) experiments have shown that the N1 component (i.e., a component related to familiar written language processing) emerges right-lateralized in preschoolers at the start of reading training (e.g., Maurer *et al.*, 2006). Crucially, this early emergence of the N1 effect has been associated with letter knowledge and it also likely reflects visual familiarity with print (Maurer, Brem, Bucher, & Brandeis, 2005). The development of the characteristic left-lateralization of the component is assumed to occur via the automatization of orthographic-phonological mappings established during learning-to-read (Maurer & McCandliss, 2007; McCandliss & Noble, 2003; Maurer *et al.*, 2006; Posner & McCandliss, 2000; see also Maurer *et al.*, 2010, for evidence with adults learning an artificial script; Brem *et al.*, 2005, for the same pattern with visual training of symbols). Thus, the increase in performance in the post-training phase, coupled with the absence of differences in location-invariant and location-specific processing between the two phases, favours the idea that these effects are associated with visual expertise with the novel scripts (i.e., a reading-related perceptual expertise; see Maurer *et al.*, 2010, for similar claims).

What we should also note is that, although both the same/different matching task and target-in-string identification task have been widely used to demonstrate orthographic effects (e.g., see Duñabeitia *et al.*, 2012; Massol *et al.*, 2013; Perea *et al.*, 2016; Scaltriti & Balota, 2013; Tydgate & Grainger, 2009), they can be performed on the basis of visual representations of the stimuli – this is the reason why these tasks can be used in both pre-training and post-training phases. As a result, some effects obtained from newly learned scripts may reflect a mixture of increased visual familiarity to the new letters together with some incipient orthographic representations (i.e., a specific to letters visual expertise, see Maurer, Brandeis, & McCandliss, 2005). Instead, skilled readers, who have already automatized the orthographic-phonological mappings, would perform these tasks not only on the basis of visual familiarity, but also on the activation of the orthographic representations of the stimuli. In other words, the patterns observed in beginner readers may rely on visual familiarization with the learned scripts, whereas the effects of skilled readers may represent an interaction between early visual processing at the letter level and feedback from orthographic representations (see Marinus *et al.*, 2018).

Taken together, our findings suggest that in order to boost the automatization of the orthographical-phonological mappings and a more parallel coarse-coding processing, a much longer learning-to-read period may be required. We now discuss several options for further research. The first option would be to run a large-scale longitudinal experiment with pre-literate children – for the sake of the argument, we assume that the experimental tasks would allow a meaningful comparison across age (see Perea *et al.*, 2016, for discussion). The experimental design would include three scripts: Roman letters, digits (i.e., a familiar visual object), and a control artificial script. This would allow us to examine not only the emergence of location-invariant and location-specific processing in a natural setting (i.e., children learning-to-read), but also how these orthographic markers vary as a consequence of literacy acquisition. Furthermore, this design would also allow examining the variations due to orthographic processing (i.e., specific to letters) vs. visual familiarity (numbers vs. artificial letters). A second option would be to train adult participants for a long period of time in an ecological setting – instead of the ecological limitations of learning an artificial

script. The most realistic design would be run the experiment with adults who are starting to learn a new language that uses an unfamiliar alphabetic script (e.g., Georgian, Armenian). In either scenario, it would be desirable to complement the behavioural tasks with the recording of brain activity during training, as this may help to disentangle the effects due to orthographical–phonological decoding from the effects due to visual training (e.g., see Maurer, Brem, *et al.*, 2005, 2006; Pleisch *et al.*, 2019, for evidence of print sensitivity in the N1 ERP amplitude; see Pleisch *et al.*, 2019, for evidence of changes in the activation of crucial orthographic processing brain regions [ventral occipitotemporal cortex and left fusiform gyrus] in the first steps of reading acquisition).

To sum up, we conducted two experiments that examined whether two markers of orthographic processing (location-invariant and location-specific processing) arise rapidly after learning-to-read and write a new script. Notably, examining when these effects emerge is essential to help interpret the subsequent developmental trajectory of orthographic effects. While participants were able to read and write with some fluency in the new script and showed some rudimentary orthographic processing, we found no evidence favouring the hypotheses that location-invariance and location-specific processing emerge quickly after learning-to-read. Instead, the emergence of these two markers of orthographic processing may take much more time, probably via the automatization of orthographic-phonological mappings.

Conflicts of interest

The authors declare no conflict of interest.

Funding

This study was supported by the Spanish Ministry of Science, Innovation, and Universities (PRE2018-083922, PSI2017-86210-P) and by the Department of Innovation, Universities, Science and Digital Society of the Valencian Government (GV/2020/074)

Author contribution

María Fernández-López (Conceptualization; Methodology; Writing – original draft) Ana Marcet (Conceptualization; Methodology; Writing – original draft) Manuel Perea, PhD (Conceptualization; Funding acquisition; Methodology; Writing – original draft)

Data Availability Statement

All the training materials, data and script files are available on the following OSF site: https://osf.io/um6rw/?view_only=7d4754bbb5f445adb5e34530162ba552

References

- Bates, D., Mächler, M., Bolker, B., & Walker, S. (2015). Fitting linear mixed-effects models using lme4. *Journal of Statistical Software*, 67, 1–48. <https://doi.org/10.18637/jss.v067.i01>
- Bates, D., Maechler, M., Bolker, B., & Walker, S. (2019). lme4: Linear Mixed-Effects Models using Eigen and S4 (1.1-20) [software]. Retrieved from <https://CRAN.R-project.org/package=lme4>.

- Brem, S., Hunkeler, E., Mächler, M., Kronschnabel, J., Karipidis, I. I., Pleisch, G., & Brandeis, D. (2018). Increasing expertise to a novel script modulates the visual N1 ERP in healthy adults. *International Journal of Behavioral Development*, 42, 333–341. <https://doi.org/10.1177/0165025417727871>
- Brem, S., Lang-Dullenkopf, A., Maurer, U., Halder, P., Bucher, K., & Brandeis, D. (2005). Neurophysiological signs of rapidly emerging visual expertise for symbol strings. *Neuroreport*, 16, 45–48. <https://doi.org/10.1097/00001756-200501190-00011>
- Bürkner, P. C. (2016). brms: An R package for Bayesian multilevel models using Stan. *Journal of Statistical Software*, 80, 1–28.
- Castles, A., Davis, C., Cavalot, P., & Forster, K. (2007). Tracking the acquisition of orthographic skills in developing readers: Masked priming effects. *Journal of Experimental Child Psychology*, 97, 165–182. <https://doi.org/10.1016/j.jecp.2007.01.006>
- Chanceaux, M., & Grainger, J. (2012). Serial position effects in the identification of letters, digits, symbols, and shapes in peripheral vision. *Acta Psychologica*, 141, 149–158. <https://doi.org/10.1016/j.actpsy.2012.08.001>
- Chanceaux, M., Mathôt, S., & Grainger, J. (2014). Effects of number, complexity, and familiarity of flankers on crowded letter identification. *Journal of Vision*, 14, 1–17. <https://doi.org/10.1167/14.6.7>
- Chetail, F. (2017). What do we do with what we learn? Statistical learning of orthographic regularities impacts written word processing. *Cognition*, 163, 103–120. <https://doi.org/10.1016/j.cognition.2017.02.015>
- R Core Team (2019). *R: A language and environment for statistical computing*. Vienna: R Foundation for Statistical Computing. Retrieved from <http://www.R-project.org/>.
- Dandurand, F., Grainger, J., & Dufau, S. (2010). Learning location-invariant orthographic representations for printed words. *Connection Science*, 22, 25–42. <https://doi.org/10.1080/09540090903085768>
- Davis, C. J. (2010). The spatial coding model of visual word identification. *Psychological Review*, 117, 713–758. <https://doi.org/10.1037/a0019738>
- Dehaene, S., Cohen, L., Sigman, M., & Vinckier, F. (2005). The neural code for written words: a proposal. *Trends in Cognitive Sciences*, 9, 335–341. <https://doi.org/10.1016/j.tics.2005.05.004>
- Duñabeitia, J. A., Dimitropoulou, M., Grainger, J., Hernández, J. A., & Carreiras, M. (2012). Differential sensitivity of letters, numbers, and symbols to character transpositions. *Journal of Cognitive Neuroscience*, 24, 1610–1624. https://doi.org/10.1162/jocn_a_00180
- Duñabeitia, J. A., Lallier, M., Paz-Alonso, P. M., & Carreiras, M. (2015). The impact of literacy on position uncertainty. *Psychological Science*, 26, 548–550. <https://doi.org/10.1177/0956797615569890>
- Duñabeitia, J. A., Orihuela, K., & Carreiras, M. (2014). Orthographic coding in illiterates and literates. *Psychological Science*, 25, 1275–1280. <https://doi.org/10.1177/0956797614531026>
- Forster, K. I., & Forster, J. C. (2003). DMDX: A Windows display program with millisecond accuracy. *Behavior Research Methods*, 35, 116–124. <https://doi.org/10.3758/BF03195503>
- García-Orza, J., Perea, M., & Muñoz, S. (2010). Are transposition effects specific to letters? *The Quarterly Journal of Experimental Psychology*, 63, 1603–1618. <https://doi.org/10.1080/17470210903474278>
- Gomez, P., Ratcliff, R., & Perea, M. (2008). The overlap model: A model of letter position coding. *Psychological Review*, 115, 577–600. <https://doi.org/10.1037/a0012667>
- Grainger, J. (2008). Cracking the orthographic code: An introduction. *Language and Cognitive Processes*, 23, 1–35. <https://doi.org/10.1080/01690960701578013>
- Grainger, J. (2018). Orthographic processing: A ‘mid-level’ vision of reading: The 44th Sir Frederic Bartlett Lecture. *Quarterly Journal of Experimental Psychology*, 71, 335–359. <https://doi.org/10.1080/17470218.2017.1314515>
- Grainger, J., Bertrand, D., Lété, B., Beyersmann, E., & Ziegler, J. C. (2016). A developmental investigation of the first-letter advantage. *Journal of Experimental Child Psychology*, 152, 161–172. <https://doi.org/10.1016/j.jecp.2016.07.016>

- Grainger, J., & Hannagan, T. (2014). What is special about orthographic processing? *Written Language & Literacy*, 17, 225–252. <https://doi.org/10.1075/wll.17.2.03gra>
- Grainger, J., Lété, B., Bertrand, D. D. S., & Ziegler, J. (2012). Evidence for multiple routes in learning to read. *Cognition*, 123, 280–292. <https://doi.org/10.1016/j.cognition.2012.01.003>
- Grainger, J., Tydgat, I., & Isselée, J. (2010). Crowding affects letters and symbols differently. *Journal of Experimental Psychology: Human Perception and Performance*, 36, 673–688. <https://doi.org/10.1037/a0016888>
- Grainger, J., & van Heuven, W. J. B. (2004). Modeling letter position coding in printed word perception. In P. Bonin (Ed.), *Mental lexicon: Some words to talk about words* (pp. 1–23). Hauppauge, NY: Nova Science Publishers.
- Grainger, J., & Ziegler, J. (2011). A dual-route approach to orthographic processing. *Frontiers in psychology*, 2, 54. <https://doi.org/10.3389/fpsyg.2011.00054>
- Krueger, L. E. (1978). A theory of perceptual matching. *Psychological Review*, 85, 278–304. <https://doi.org/10.1037//0033-295X.85.4.278>
- Lally, C., Taylor, J. S. H., Lee, C. H., & Rastle, K. (2020). Shaping the precision of letter position coding by varying properties of a writing system. *Language, Cognition and Neuroscience*, 35, 374–382. <https://doi.org/10.1080/23273798.2019.1663222>
- Leinenger, M., Myslín, M., Rayner, K., & Levy, R. (2017). Do resource constraints affect lexical processing? Evidence from eye movements. *Journal of Memory and Language*, 93, 82–103. <https://doi.org/10.1016/j.jml.2016.09.002>
- Marcet, A., Perea, M., Baciero, A., & Gomez, P. (2019). Can letter position encoding be modified by visual perceptual elements? *Quarterly Journal of Experimental Psychology*, 72, 1344–1353. <https://doi.org/10.1177/1747021818789876>
- Marinus, E., Kezilas, Y., Kohnen, S., Robidoux, S., & Castles, A. (2018). Who are the noisiest neighbors in the hood? Using error analyses to study the acquisition of letter-position processing. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 44, 1384–1396. <https://doi.org/10.1037/xlm0000524>
- Massol, S., Duñabeitia, J. A., Carreiras, M., & Grainger, J. (2013). Evidence for letter-specific position coding mechanisms. *PLoS One*, 8, e68460. <https://doi.org/10.1371/journal.pone.0068460>
- Maurer, U., Blau, V. C., Yoncheva, Y. N., & McCandliss, B. D. (2010). Development of visual expertise for reading: rapid emergence of visual familiarity for an artificial script. *Developmental Neuropsychology*, 35, 404–422. <https://doi.org/10.1080/87565641.2010.480916>
- Maurer, U., Brandeis, D., & McCandliss, B. D. (2005). Fast, visual specialization for reading in English revealed by the topography of the N170 ERP response. *Behavioral and Brain Functions*, 1, 13. <https://doi.org/10.1186/1744-9081-1-13>
- Maurer, U., Brem, S., Bucher, K., & Brandeis, D. (2005). Emerging neurophysiological specialization for letter strings. *Journal of Cognitive Neuroscience*, 17, 1532–1552. <https://doi.org/10.1162/089892905774597218>
- Maurer, U., Brem, S., Kranz, F., Bucher, K., Benz, R., Halder, P., . . . Brandeis, D. (2006). Coarse neural tuning for print peaks when children learn to read. *NeuroImage*, 33, 749–758. <https://doi.org/10.1016/j.neuroimage.2006.06.025>
- Maurer, U., & McCandliss, B. D. (2007). The development of visual expertise for words: the contribution of electrophysiology. In E. L. Grigorenko & A. J. Naples (Eds.), *Single-word reading: cognitive, behavioral and biological perspectives* (pp. 43–64). Mahwah, NJ: Lawrence Erlbaum Associates.
- McCandliss, B. D., & Noble, K. G. (2003). The development of reading impairment: a cognitive neuroscience model. *Mental retardation and developmental disabilities research reviews*, 9, 196–205. <https://doi.org/10.1002/mrdd.10080>
- Morey, R. D., & Rouder, J. N. (2018). Package ‘BayesFactor’. R Package version 0.9.12-4.2. Retrieved from <https://richarddmorey.github.io/BayesFactor/>.
- Muñoz, S., Perea, M., García-Orza, J., & Barber, H. A. (2012). Electrophysiological signatures of masked transposition priming in a same-different task: Evidence with strings of letters vs. pseudoletters. *Neuroscience Letters*, 515, 71–76. <https://doi.org/10.1016/j.neulet.2012.03.021>

- Norris, D., & Kinoshita, S. (2012). Reading through a noisy channel: Why there's nothing special about the perception of orthography. *Psychological Review*, 119, 517–545. <https://doi.org/10.1037/a0028450>
- Norris, D., Kinoshita, S., & van Casteren, M. (2010). A stimulus sampling theory of letter identity and order. *Journal of Memory and Language*, 62, 254–271. <https://doi.org/10.1016/j.jml.2009.11.002>
- Perea, M., Jiménez, M., & Gomez, P. (2016). Does location uncertainty in letter position coding emerge because of literacy training? *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 42, 996–1001. <https://doi.org/10.1037/xlm0000208>
- Perea, M., & Lupker, S. J. (2004). Can CANISO activate CASINO? Transposed-letter similarity effects with nonadjacent letter positions. *Journal of Memory and Language*, 51, 231–246. <https://doi.org/10.1016/j.jml.2004.05.005>
- Pleisch, G., Karipidis, I. I., Brauchli, C., Röthlisberger, M., Hofstetter, C., Stämpfli, P., . . . Brem, S. (2019). Emerging neural specialization of the ventral occipitotemporal cortex to characters through phonological association learning in preschool children. *NeuroImage*, 189, 813–831. <https://doi.org/10.1016/j.neuroimage.2019.01.046>
- Posner, M., & McCandliss, B. D. (2000). Brain circuitry during reading. In R. Klein & P. McMullen (Eds.), *Converging methods for understanding reading and dyslexia* (pp. 305–337). Cambridge: MIT Press.
- Price, C. J., & Devlin, J. T. (2011). The interactive account of ventral occipitotemporal contributions to reading. *Trends in Cognitive Sciences*, 15, 246–253. <https://doi.org/10.1016/j.tics.2011.04.001>
- Ratcliff, R. (1981). A theory of order relations in perceptual matching. *Psychological Review*, 88, 552–572. <https://doi.org/10.1037/0033-295X.88.6.552>
- Reicher, G. M. (1969). Perceptual recognition as a function of meaningfulness of stimulus material. *Journal of Experimental Psychology*, 81, 275–280. <https://doi.org/10.1037/h0027768>
- Rouder, J. N., Speckman, P. L., Sun, D., Morey, R. D., & Iverson, G. (2009). Bayesian t tests for accepting and rejecting the null hypothesis. *Psychonomic Bulletin & Review*, 16, 225–237. <https://doi.org/10.3758/PBR.16.2.225>
- Scaltritti, M., & Balota, D. A. (2013). Are all letters really processed equally and in parallel? Further evidence of a robust first letter advantage. *Acta Psychologica*, 144, 397–410. <https://doi.org/10.1016/j.actpsy.2013.07.018>
- Scaltritti, M., Dufau, S., & Grainger, J. (2018). Stimulus orientation and the first-letter advantage. *Acta Psychologica*, 183, 37–42. <https://doi.org/10.1016/j.actpsy.2017.12.009>
- Schönbrodt, F. D., & Wagenmakers, E. J. (2018). Bayes factor design analysis: Planning for compelling evidence. *Psychonomic Bulletin & Review*, 25, 128–142. <https://doi.org/10.3758/s13423-017-1230-y>
- Shubert, T., Badcock, N., & Kohnen, S. (2017). Development of children's identity and position processing for letter, digit, and symbol strings: A cross-sectional study of the primary school years. *Journal of Experimental Child Psychology*, 162, 163–180. <https://doi.org/10.1016/j.jecp.2017.05.008>
- Staub, A., & Goddard, K. (2019). The role of preview validity in predictability and frequency effects on eye movements in reading. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 45, 110–127. <https://doi.org/10.1037/xlm0000561>
- Taylor, J. S. H., Davis, M. H., & Rastle, K. (2017). Comparing and validating methods of reading instruction using behavioural and neural findings in an artificial orthography. *Journal of Experimental Psychology: General*, 146, 826–858. <https://doi.org/10.1037/xge0000301>
- Taylor, J. S. H., Plunkett, K., & Nation, K. (2011). The influence of consistency, frequency, and semantics on learning to read: An artificial orthography paradigm. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 37, 60–76. <https://doi.org/10.1037/a0020126>
- Tydgat, I., & Grainger, J. (2009). Serial position effects in the identification of letters, digits, and symbols. *Journal of Experimental Psychology: Human Perception and Performance*, 35, 480–498. <https://doi.org/10.1037/a0013027>

- Vejnović, D., & Zdravković, S. (2015). Side flankers produce less crowding, but only for letters. *Cognition*, 143, 217–227. <https://doi.org/10.1016/j.cognition.2015.07.003>
- Vidal, C., & Chetail, F. (2017). BACS: The Brussels Artificial Character Sets for studies in cognitive psychology and neuroscience. *Behavior Research Methods*, 49, 2093–2112. <https://doi.org/10.3758/s13428-016-0844-8>
- Wagenmakers, E. J., Love, J., Marsman, M., Jamil, T., Ly, A., Verhagen, J., ... Meerhoff, F. (2018). Bayesian inference for psychology. Part II: Example applications with JASP. *Psychonomic Bulletin & Review*, 25, 58–76. <https://doi.org/10.3758/s13423-017-1323-7>
- Wagenmakers, E.-J., Marsman, M., Jamil, T., Ly, A., Verhagen, J., Love, J., ... Morey, R. D. (2017). Bayesian inference for psychology. Part I: Theoretical advantages and practical ramifications. *Psychonomic Bulletin & Review*, 25, 35–57. <https://doi.org/10.3758/s13423-017-1343-3>
- Wheeler, D. D. (1970). Processes in word recognition. *Cognitive Psychology*, 1, 59–85. [https://doi.org/10.1016/0010-0285\(70\)90005-8](https://doi.org/10.1016/0010-0285(70)90005-8)
- Whitney, C. (2001). How the brain encodes the order of letters in a printed word: The SERIOL model and selective literature review. *Psychonomic Bulletin and Review*, 8, 221–243. <https://doi.org/10.3758/BF03196158>
- Winkel, H., Perea, M., & Peart, E. (2014). Testing the flexibility of the modified receptive field (MRF) theory: Evidence from an unspaced orthography (Thai). *Acta Psychologica*, 150, 55–60. <https://doi.org/10.1016/j.actpsy.2014.04.008>
- Ziegler, J. C., Bertrand, D., Lété, B., & Grainger, J. (2014). Orthographic and phonological contributions to reading development: Tracking developmental trajectories using masked priming. *Developmental Psychology*, 50, 1026–1036. <https://doi.org/10.1037/a0035187>
- Ziegler, J. C., Pech-Georgel, C., Dufau, S., & Grainger, J. (2010). Rapid processing of letters, digits and symbols: what purely visual-attentional deficit in developmental dyslexia? *Developmental Science*, 13, F8–F14. <https://doi.org/10.1111/j.1467-7687.2010.00983.x>

Received 26 September 2019; revised version received 18 June 2020

Appendix A:

Supplementary analyses with Bayesian linear mixed-effects models

Experiment 1: location-invariant processing

For the sake of completeness, we also examined the latency and accuracy data of the confirmatory analyses using Bayesian linear mixed-effects models using the brms package in R (Bürkner, 2016). An advantage of this procedure – via Stan – is that it allows us to fit the models using the maximal random-effect structure (see Bates, Mächler, Bolker & Walker, 2015, for arguments in favour of maximal models).

Different trials

Trained script vs. Untrained script. The fitted model was: Dependent_Variable [i.e., –1000/RT or accuracy] ~ script * condition * phase + (1 + script * condition * phase | subject) + (1 + condition * phase | item). More complex random-effects terms resulted in model non-convergence. Furthermore, these models offer the Bayesian 95% credible intervals for each parameter based on the posterior distributions. For the latency data, we employed the same response time transformation as in LME analyses (i.e., –1000/RT; family = gaussian), whereas for the accuracy data, we used the Bernoulli distribution (family = bernoulli) – this is the parallel to family = binomial in GLME models. We

Table A1. Parameter estimates in latency and accuracy supplemental analyses of Experiment I ('different' trials: trained vs. untrained script)

Latency data	Estimate	SE	95% Credible Interval
Intercept	-1.84	0.06	[-1.95, -1.73]
Script	-0.01	0.02	[-0.05, 0.02]
Condition	0.08	0.01	[0.05, 0.10]
Phase	0.13	0.04	[0.06, 0.21]
Script × Condition	-0.01	0.02	[-0.05, 0.03]
Script × Phase	-0.05	0.04	[-0.12, 0.03]
Condition × Phase	-0.01	0.02	[-0.05, 0.03]
Script × Condition × Phase	0.01	0.03	[-0.04, 0.06]
Accuracy data			
Intercept	1.56	0.15	[1.27, 1.87]
Script	-0.07	0.10	[-0.28, 0.13]
Condition	-1.06	0.11	[-1.28, -0.84]
Phase	-0.35	0.12	[-0.60, -0.10]
Script × Condition	0.07	0.12	[-0.17, 0.31]
Script × Phase	-0.04	0.13	[-0.29, 0.21]
Condition × Phase	0.12	0.13	[-0.13, 0.37]
Script × Condition × Phase	0.03	0.17	[-0.29, 0.36]

Note. Those effects with 95% credible intervals beyond 0 are in bold.

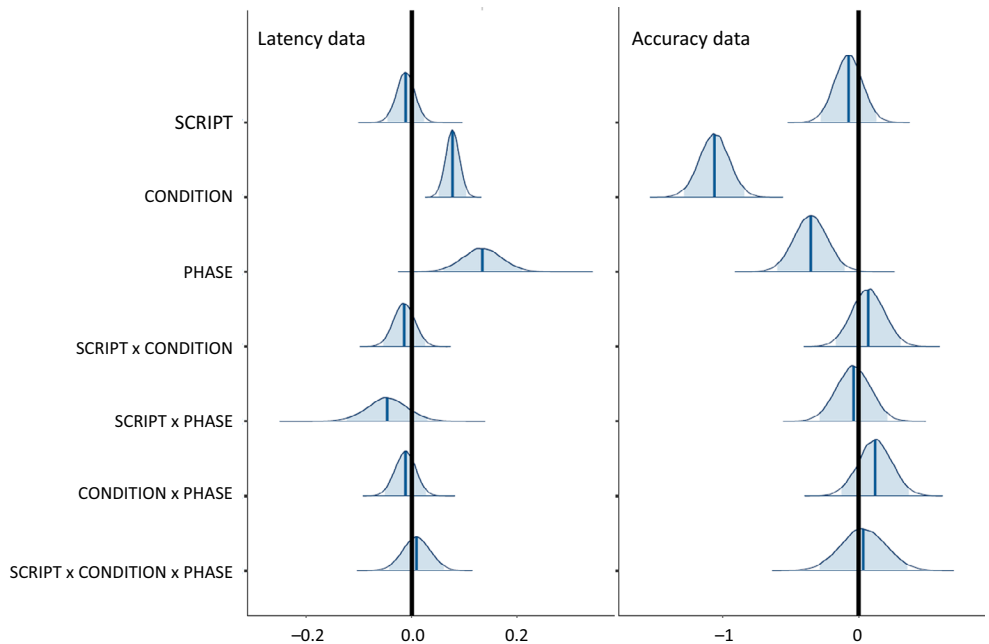


Figure A1. Posterior effects estimates from the Bayesian linear mixed models for 'different' trials in Experiment I (Trained vs. Untrained script) (left panel: latency analysis; right panel: accuracy analysis). The thick black line corresponds to an effect of zero, the dark grey line corresponds to the estimates, and the shaded area corresponds to the 95% credible interval. [Colour figure can be viewed at wileyonlinelibrary.com]

Table A2. Parameter estimates in latency and accuracy supplemental exploratory analyses of Experiment I ('different' trials: Roman vs. trained script)

Latency data	Estimate	SE	95% Credible Interval
Intercept	−1.83	0.04	[−1.91, −1.74]
Condition	0.07	0.02	[0.04, 0.11]
Script	−0.02	0.05	[−0.13, 0.09]
Condition × Script	0.00	0.02	[−0.04, 0.05]
Accuracy data			
Intercept	1.86	0.17	[1.55, 2.20]
Condition	−1.66	0.14	[−1.92, −1.39]
Script	−0.26	0.21	[−0.69, 0.16]
Condition × Script	0.56	0.17	[0.23, 0.90]

Note. Those effects with 95% credible intervals beyond 0 are in bold.

employed 4 chains, each with 10,000 iterations after a warm-up of 1000 iterations. The maximal random-effect structure models converged successfully: the values of Rhat were 1.00 for all parameters.

As can be seen from the estimates and 95% credible intervals presented in Table A1, we found robust evidence of an effect of Condition (i.e., probe-target relationship) and Phase in both latency and accuracy analyses, thus corroborating the pre-registered

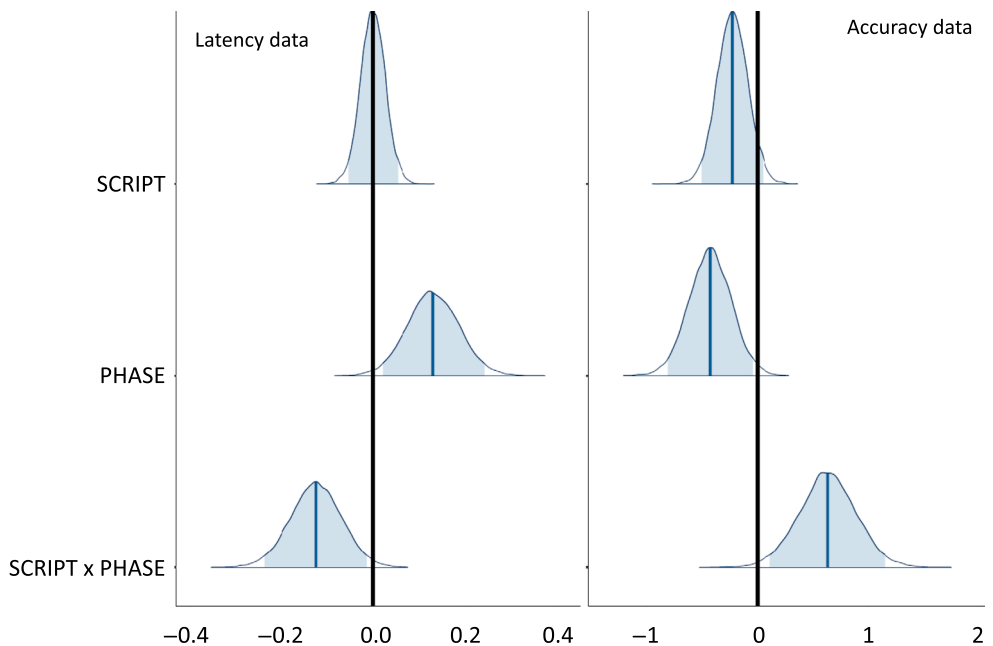


Figure A2. Posterior effects estimates from the Bayesian linear mixed models in 'different' trials in Experiment I (Roman vs. Trained script) (left panel: latency analysis; right panel: accuracy analysis). The thick black line corresponds to an effect of zero, the dark grey line corresponds to the estimate, and the shaded area corresponds to the 95% credible interval. [Colour figure can be viewed at wileyonlinelibrary.com]

analyses. Similarly, we did not find any evidence of a three-way interaction between Script, Condition, and Phase (see Figure A1, for the posterior distributions).

Roman script vs. trained script

The fixed factors were Script (Trained [post-training] vs. Roman) and Condition (transposed, replaced). We followed the same procedure as above, with 5,000 iterations (Rhat = 1.00 in all cases). Table A2 shows the estimates and 95% credible intervals from the latency and accuracy models, and Figure A2 shows the posterior distributions of the parameters. Together with a substantial transposed-letter effect in the latency data, these analyses confirmed the interaction between Script and Condition in accuracy data.

Same trials

Trained script vs. Untrained script. The general procedure was the same as above (i.e., the maximal random-effect structure model with 5,000 iterations; Rhat = 1.00 in all cases). As shown in Table A3 (estimates and 95% credible intervals) and Figure A3 (posterior distributions of the parameters), these analyses corroborated the interaction effect for between Phase and Script in the latency and accuracy data.

Experiment 2: location-specific processing

Trained vs. Untrained script

As in Experiment 1, we examined the data with Bayesian linear mixed-effects models using the brms package (Bürkner, 2016) in R. The three fixed factors were the same as in the pre-registered analyses. The initial letter was set as the reference level for the factor Position. We fitted the maximal random-effect structure model (i.e., $\text{accuracy} \sim \text{script} * \text{position} * \text{phase} + (1 + \text{script} * \text{position} * \text{phase} | \text{subject}) + (1 + \text{position} * \text{phase} | \text{item})$) using 4 chains, each with 5,000 iterations after a warm-up of 1,000 iterations. The priors were the same as in Experiment 1. The model converged successfully (Rhat = 1.00 for all parameters).

As can be seen in Table A4, accuracy in the initial-letter position was substantially lower than in the middle, fixated position. The accuracy advantage of the middle position

Table A3. Parameter estimates in latency and accuracy supplemental analyses of Experiment 1 ('same' trials: trained vs. untrained script)

Latency data	Estimate	SE	95% Credible Interval
Intercept	-1.99	0.08	[-2.14, -1.84]
Script	0.00	0.03	[-0.05, 0.05]
Phase	0.13	0.06	[0.02, 0.24]
Script × Phase	-0.13	0.06	[-0.24, -0.01]
Accuracy data			
Intercept	2.70	0.21	[2.31, 3.12]
Script	-0.23	0.14	[-0.51, 0.05]
Phase	-0.43	0.19	[-0.81, -0.05]
Script × Phase	0.64	0.26	[0.11, 1.15]

Note. Those effects with 95% credible intervals beyond 0 are in bold.

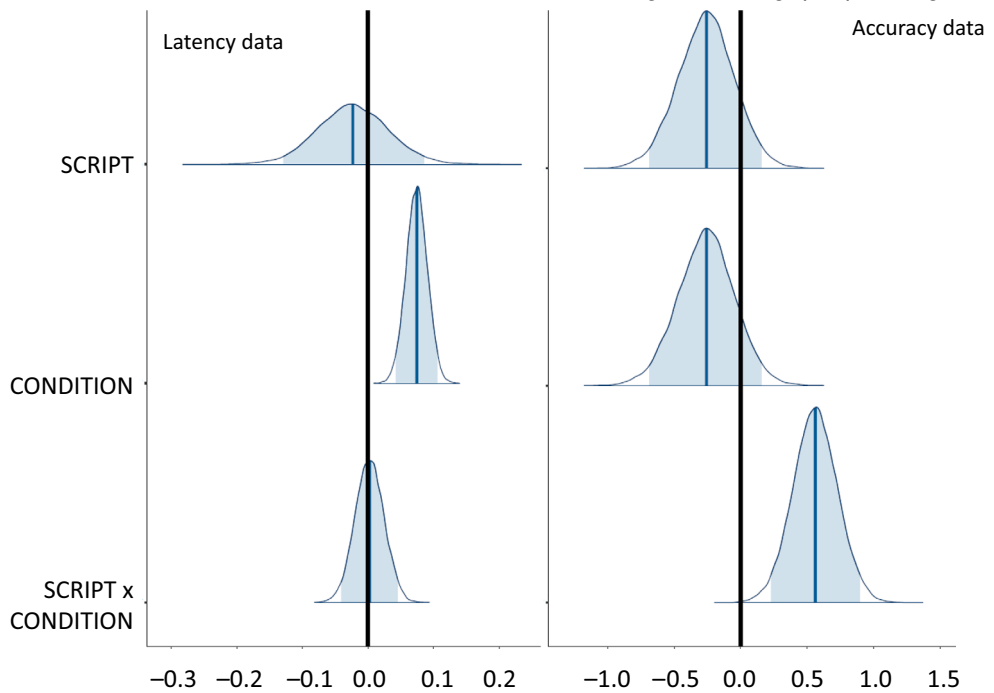


Figure A3. Posterior effects estimates from the Bayesian linear mixed models in 'same' trials of Experiment 1. The thick black line corresponds to an effect of zero, the dark grey line corresponds to the estimates, and the shaded area corresponds to the 95% credible interval. [Colour figure can be viewed at wileyonlinelibrary.com]

Table A4. Parameter estimates in the accuracy supplemental analyses of Experiment 2 (trained vs. untrained script)

		Estimate	SE	95% Credible Interval
Intercept		0.16	0.09	[−0.03, 0.34]
Script		−0.05	0.13	[−0.29, 0.20]
Position	2 nd	−0.22	0.14	[−0.50, 0.06]
	3 rd	1.89	0.26	[1.38, 2.41]
	4 th	−0.20	0.14	[−0.48, 0.07]
	5 th	−0.14	0.14	[−0.41, 0.13]
Phase		0.04	0.14	[−0.22, 0.31]
Script x	2 nd position	0.07	0.18	[−0.29, 0.42]
	3 rd position	−0.36	0.22	[−0.80, 0.08]
	4 th position	0.13	0.18	[−0.22, 0.49]
	5 th position	−0.03	0.18	[−0.39, 0.33]
Script × Phase		−0.06	0.18	[−0.41, 0.29]
Phase ×	2 nd position	0.07	0.20	[−0.33, 0.46]
	3 rd position	−0.80	0.29	[−1.38, −0.22]
	4 th position	−0.03	0.20	[−0.42, 0.37]
	5 th position	0.09	0.20	[−0.29, 0.49]
Script × Phase ×	2 nd position	0.08	0.25	[−0.41, 0.58]
	3 rd position	0.23	0.31	[−0.38, 0.83]
	4 th position	0.21	0.25	[−0.28, 0.71]
	5 th position	−0.02	0.26	[−0.52, 0.48]

Note. Those effects with 95% credible intervals beyond 0 are in bold.

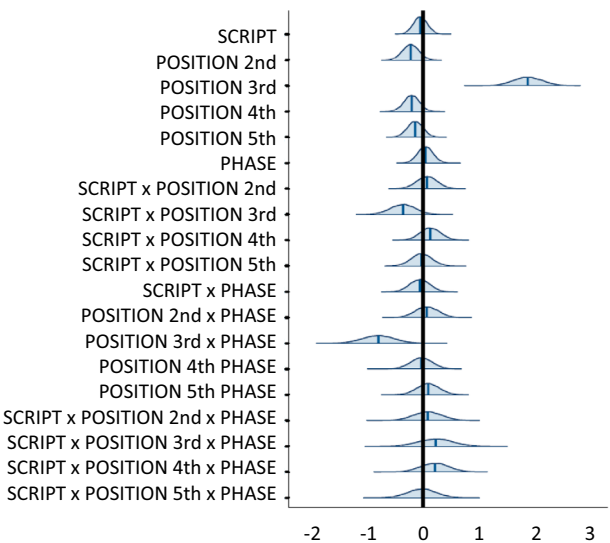


Figure A4. Posterior effects estimates from the Bayesian linear mixed models in Experiment 2. The thick black line corresponds to an effect of zero, the dark grey line corresponds to the estimate, and the shaded area corresponds to the 95% credible interval. [Colour figure can be viewed at wileyonlinelibrary.com]

increased in the post-trained phase, as deduced from the interaction with Phase. In addition, the parameter estimates showed some advantage of the initial position over the other letter positions (see Figure A4, for the posterior effects). Therefore, these analyses corroborate the findings obtained in the ANOVAs indicated in the pre-registered analyses.

Roman script vs. learned script. The procedure was the same as above, except that the two fixed factors were Script (Roman vs. Trained) and Position – the reference level for the factor Position was also the first letter. The maximal model converged successfully

Table A5. Parameter estimates in the accuracy supplemental analyses of Experiment 2 (Roman script vs. learned script)

		Estimate	SE	95% Credible Interval
Intercept		0.15	0.10	[−0.04, 0.35]
Script		0.01	0.14	[−0.25, 0.29]
Position	2 nd	−0.22	0.15	[−0.51, 0.06]
	3 rd	1.90	0.28	[1.38, 2.47]
	4 th	−0.20	0.14	[−0.49, 0.08]
	5 th	−0.14	0.15	[−0.43, 0.15]
Script x	2 nd position	0.15	0.21	[−0.25, 0.55]
	3 rd position	−0.14	0.33	[−0.77, 0.53]
	4 th position	0.29	0.20	[−0.11, 0.69]
	5 th position	0.18	0.21	[−0.24, 0.59]

Note. Those effects with 95% credible intervals beyond 0 are in bold.

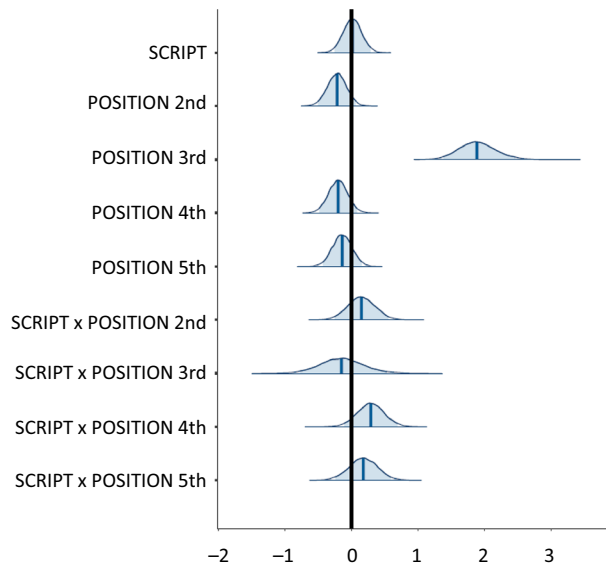


Figure A5. Posterior effects estimates from the Bayesian linear mixed models in Experiment 2 (Roman vs. Trained script). The thick black line corresponds to an effect of zero, the dark grey line corresponds to the estimate, and the shaded area corresponds to the 95% credible interval. [Colour figure can be viewed at wileyonlinelibrary.com]

($R^2 = 1.00$ for all parameters). As can be seen from the 95% credible intervals (see Table A5; see also Figure A5, for the posterior distributions), we found a substantial advantage of the third letter position. In addition, there was a numerical advantage of the first-letter position relative to the second letter position – this effect did not interact with script.

Appendix B:

Figures B1 and B2 provide a visual representation of how training improved participants' performance along the learning days (from day 2 to day 6 – note that, on day 1, participants only had to listen to and write down the phoneme–grapheme correspondences).

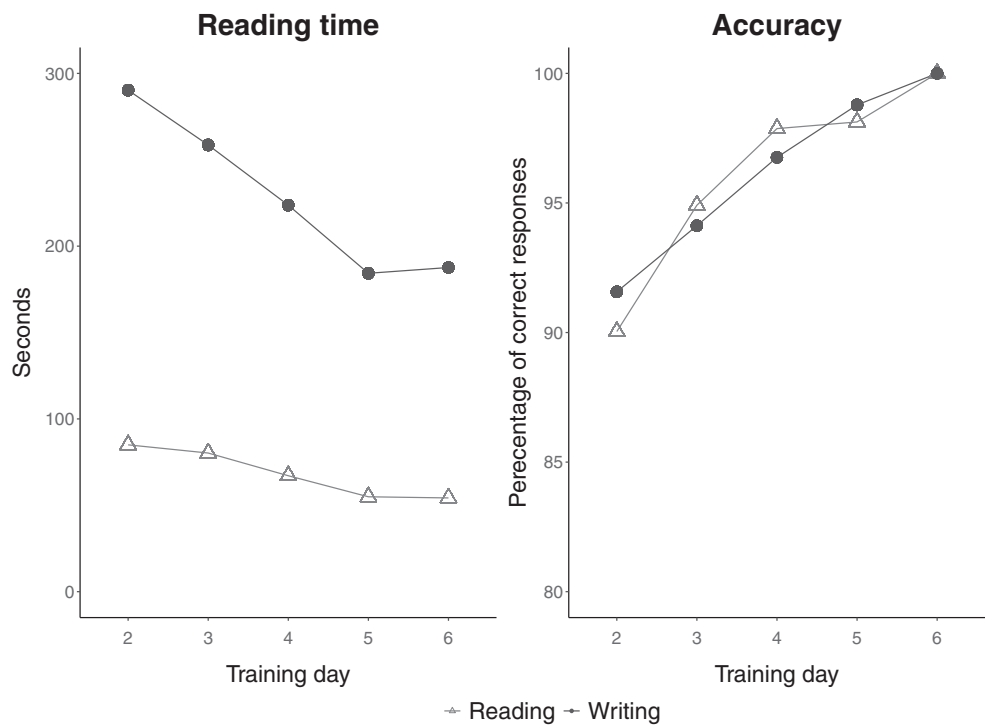


Figure B1. Participants performance on the reading aloud and writing tasks along the training days (from Day 2 to Day 6).

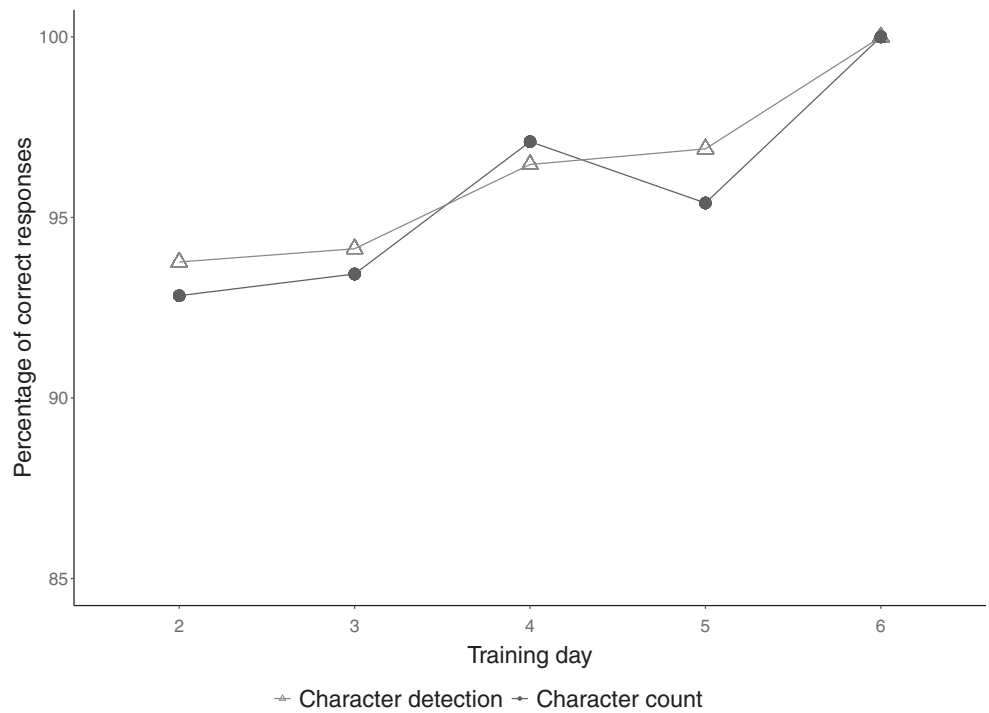


Figure B2. Participants performance on the visual familiarization tasks along the training days (from Day 2 to Day 6).

How Resilient Is Reading to Letter Rotations? A Parafoveal Preview Investigation

María Fernández-López¹, Jonathan Mirault², Jonathan Grainger², and Manuel Perea^{1, 3}

¹ Departamento de Metodología and ERI-Lectura, Universitat de València

² Laboratoire de Psychologie Cognitive, Centre National de la Recherche Scientifique and Aix-Marseille Université

³ Centro de Ciencia Cognitiva, Universidad Nebrija

Skilled readers have developed a certain amount of tolerance to variations in the visual form of words (e.g., CAPTCHAs, handwritten text, etc.). To examine how visual distortion affects the mapping from the visual input onto abstract word representations during normal reading, we focused on a single type of distortion: letter rotation. Importantly, two leading neurally inspired models of word recognition (SERIOL model, Whitney, 2001; LCD model, Dehaene et al., 2005) make distinct predictions: Whereas the SERIOL model postulates that the cost of letter rotation increases gradually, the LCD model assumes an all-or-none boundary at around 40° to 45° rotation angle. To examine these predictions in a normal reading scenario, we conducted a parafoveal preview experiment using the gaze-contingent boundary change paradigm. The parafoveal previews were identical or unrelated to the target word. Critically, each preview's individual letters were rotated 15°, 30°, 45°, or 60°. Apart from the parafoveal previews, all text, including target words, was presented with letters in their canonical upright orientation. Results showed that the advantage of the identity preview condition in eye fixation times on the target word decreased progressively by rotation angle. This pattern of results favors the view that the cost of letter rotation during normal reading increases gradually as a function of the angle of rotation.

Keywords: eye-movements, letter rotation, parafoveal processing, reading

Efficient sentence reading requires the rapid abstraction of each word's visual form onto multiple levels of information (orthography, phonology, semantics, and syntax) while the eyes move along the text (see Clifton et al., 2016; Rayner, 2009; for reviews). Because of their extensive experience across different reading conditions (e.g., numerous fonts; uppercase TEXT, handwritten text, cursive, CAPTCHAs, etc.), skilled readers have developed some tolerance to noise and variations in shape (i.e., visual

characteristics) of the words' constituent letters (see Grainger & Dufau, 2012).

Although often overlooked in reading research, the study of the influence of visual factors in reading and how skilled readers tolerate various kinds of shape distortion is essential for a complete understanding of the reading process (Slattery, 2016; see also Tinker, 1958; Tinker & Paterson, 1931; for reviews of early research). One of the reasons of the shortage of studies on this issue is that it is difficult to separate the effect of all the relevant variables underlying visual distortion. For instance, handwritten text, and CAPTCHAs are distorted in several dimensions, and hence any processing differences relative to the pristine written format we typically encounter when reading might be attributable to multiple factors.

To overcome this interpretative issue, in the present article we focus on a single type of visual distortion during sentence reading: the rotation of individual letters in words. An added value of studying letter rotations is that leading neural models of visual word recognition make explicit predictions about how letter rotations affect word processing. The SERIOL model of word recognition (Whitney, 2001, 2002) assumes that "input levels to letter units are reduced for rotated stimuli" (p. 119, Whitney, 2002), increasing gradually the time required to encode a letter string. Instead, the Local Combination Detectors model (henceforth, LCD model; Dehaene et al., 2005) offers an all-or-none account on how letter rotation affects letter/word processing. Specifically, the LCD model postulates that letter rotation impedes the normal

This article was published Online First April 15, 2021.

María Fernández-López <https://orcid.org/0000-0001-7419-7369>

Jonathan Mirault <https://orcid.org/0000-0003-1327-7861>

Jonathan Grainger <https://orcid.org/0000-0002-6737-5646>

Manuel Perea <https://orcid.org/0000-0002-3291-1365>

This study was supported by the Spanish Ministry of Science, Innovation, and Universities (PRE2018-083922, PSI2017-86210-P), the Department of Innovation, Universities, Science, and Digital Society of the Valencian Government (GV/2020/074), and the European Research Council (Grant ERC 742141). We thank Romain Torras for his help in running the laboratory experiment.

All the data, stimuli, and R code are available at osf.io/h7uvj.

Correspondence concerning this article should be addressed to Jonathan Mirault, Laboratoire de Psychologie Cognitive, Aix-Marseille Université, 3 place Victor Hugo, 13003 Marseille, France. Email: jonathan.mirault@univ-amu.fr

mapping of the letter features from the visual word form onto the letter detectors for angles above 40° to 45°, but that there would be little cost for rotations below that boundary¹ (see also Cohen et al., 2008; Vinckier et al., 2006). Consistent with this latter prediction, a number of experiments have shown a cost in processing when words are rotated as a whole at 45° of rotation or larger. For instance, Cohen et al. (2008) found no differences in response times to words presented with 0° and 22.5° rotation angles, whereas there was a dramatic slow-down at angles of 45° and larger. In addition, other experiments have reported slower word identification times for rotated words at angles of 45° or larger than for their horizontally-presented counterparts (e.g., Barnhart & Goldinger, 2013 [0° vs. 90°]; Gomez & Perea, 2014 [0° vs 45° vs. 90°]; Koriat & Norman, 1984 [0° vs. 60° vs. 120° vs. 180°]).

Clearly, the above-cited studies consistently showed a cost when reading rotated words at 45° or larger. However, they were not informative as to the locus of the effect. These effects could have occurred at an early encoding stage—as the LCD or SERIOL models would predict—but they could also have been attributable to a processing cost at a late verification stage of processing where the visual percept is checked against the stored lexical unit. To directly test whether the locus of the effect of rotation occurs early in word processing, Perea et al. (2018) conducted a masked priming lexical decision experiment with 90° rotation of the whole word. The rationale of their experiment was that if the encoding of orthographic-lexical representations were hindered by rotations of 45° or larger, one would expect a dramatically reduced masked priming effect for 90° rotated words. However, they found sizable masked identity and transposed-letter priming effects in line with the effects obtained with horizontally presented stimuli (see also Yang & Lupker, 2019; for evidence of robust masked priming effects with 90° and 180° rotated words). Hence, at least when words are presented in isolation, word rotation does not seem to hinder initial access to the orthographic representations (see also Perea et al., 2020, for similar evidence with isolated letters).

The robustness of masked priming effects to word rotations is quite revealing. However, one could argue that readers might rotate the whole stimulus and then process each letter as usual (see Gomez & Perea, 2014, and Whitney, 2002, for discussion). Thus, a more stringent test for the LCD and SERIOL models would be to examine the effect of rotating the individual letters within words rather than rotating the whole word. Kim and Straková (2012) conducted a lexical decision experiment in which individual letters of words and pseudowords were rotated (0°, 22.5°, 45°, 67.5°, or 90°). Importantly, they also recorded event-related potentials (ERPs) during the task. Kim and Straková (2012) found an increase in amplitude in the P1 and N170 components (i.e., two indexes that reflect the initial contact with word-form features) for all rotation angles (i.e., including 22.5°) when compared with the horizontal format. To explain this pattern, they suggested that words whose individual letters are rotated require additional resources that are recruited very early in processing. The behavioral data showed a different picture, though. For words, response times were remarkably similar for rotation angles between 0° and 67.5°, with a large increase from 67.5° to 90°. Accuracy rates were similar for angles from 0° to 45°, slightly decreased from 45° to 67.5°, and there was a large drop from 67.5° to 90°. Thus, the behavioral data from Kim and Straková (2012) suggest that word recognition was not severely affected by rotations smaller than 45°

to 67.5°—note that evidence for early ERP effects does not necessarily imply longer word identification times (see Perea et al., 2015, for discussion).

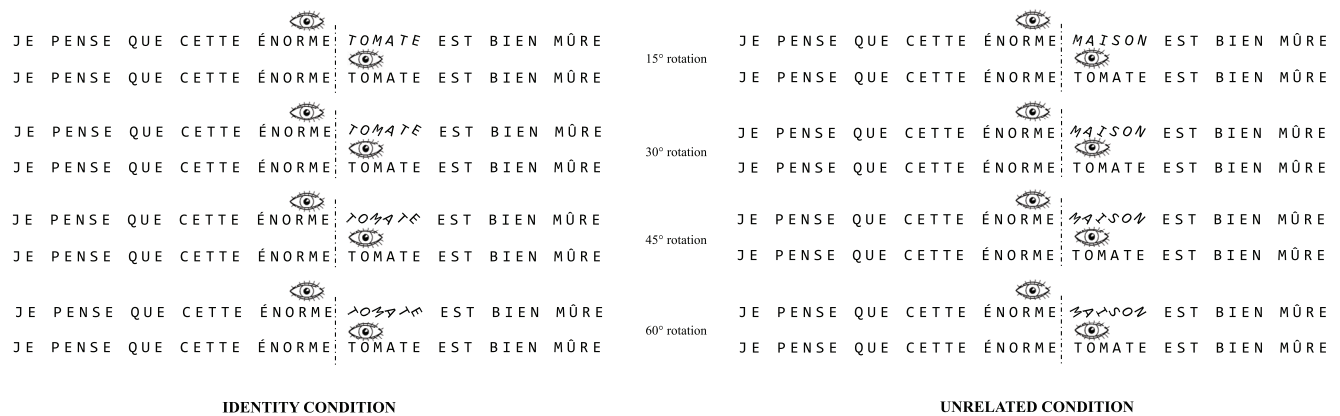
More important for the present purposes, a recent study measured participants' eye movements when reading sentences in which the individual letters within words were rotated (30°, 60°) or not (Blythe et al., 2019). Blythe et al. found an overall reading cost on sentence reading (total viewing time, number of fixations) for both rotation angles (30° and 60°) when comparing with upright presented text, and this cost was larger at the 60° than at the 30° rotation angle. They also examined whether rotation angle interacted with word frequency by embedding a high- versus low-frequency target word in each sentence. Whereas both rotation angle and word-frequency had additive effects in the first-fixation duration on the target word (30° < 60°; high-frequency < low-frequency), total fixation durations (i.e., the sum of all fixations in the target word) showed an interaction between the two factors: The magnitude of the word-frequency effect was a function of the rotation angle (i.e., the difference in eye fixation durations between high- and low-frequency target words was greater at 60°, it decreased at 30° and it was even smaller at 0°), thus suggesting that lexical access was impaired even by relatively small rotations (i.e., 30°). Blythe et al. (2019) concluded that letter detectors are affected by rotations well below 45° during sentence reading, and thereby the LCD predictions concerning letter rotations should be revised (i.e., the gradual cost of letter rotation may fit better with the SERIOL model). Although Blythe et al.'s (2019) experiment confirmed that letter rotations do affect the reading process, it does not inform us as to whether the effect of letter rotation occurs early in processing (i.e., as posited by the LCD and SERIOL models) or whether it occurs at late integration processes (keep in mind that participants could have adjusted their eye movement pattern to the rotation angle of the sentences).

The main aim of the present experiment was to examine whether the cost due to letter rotation in normal silent reading emerges early in processing, as predicted by the LCD and SERIOL models. To obtain a strict test of the resilience of letter detectors to rotations, we rotated the individual letters within words in a parafoveal word during sentence reading with the gaze-contingent boundary change paradigm (Rayner, 1975). Rayner's gaze-contingent boundary change technique allows for the manipulation of parafoveal information that is available to the reader before their eyes move to that location hence enabling foveal processing of a target word presented at the same location (see Figure 1). The rationale of this paradigm is that when we read, we extract information not only from the fixated word, but also from the following word(s) in the parafovea (Rayner et al., 1981; Rayner et al., 1982; see also Vasilev et al., 2018; for evidence with degraded previews). Thus, it is possible to infer the amount of useful information the reader obtains from the parafovea by examining how the fixation durations on a target word vary as a function of the type of parafoveal preview (e.g., identity previews vs. unrelated previews; phonological

¹ Note that the prediction of the LCD model (Dehaene et al., 2005) "letter detectors should be disrupted by rotation (>40°)" (p. 340) was based on a study conducted by Logothetis and Pauls (1995) with monkeys. Nevertheless, the empirical evidence provided by Cohen et al. (2008) and Vinckier et al. (2006) with adult readers was interpreted in light of this prediction of the model.

Figure 1

Description of an Eye Movement–Contingent Display–Change Trial With the Eight Experimental Conditions (Identity Preview Rotated at 15°, 30°, 45°, and 60°; “I Think This Huge Tomato Is Ripe”; Unrelated Preview Rotated at 15°, 30°, 45° and 60°; “I Think This Huge House Is Ripe”)



Note. The eye symbol represents where the reader is fixating, and the vertical dashed line represents the invisible boundary preceding the target word. Before crossing the boundary, the sentence is presented with the rotated previews. When the eyes cross the boundary, the parafoveal preview is replaced by the target word in the canonical format.

preview vs. orthographic control preview, etc.; see Rayner et al., 2012; for review).

In the present gaze-contingent boundary change paradigm experiment, the parafoveal preview could be the same as the target (identity preview-target pairs) or a different word (unrelated preview-target pairs). We chose identity previews because they provide the largest effects in the parafovea (see Angele et al., 2013). The letters of the identity/unrelated previews could be rotated at four different angles, the difference always being 15° (15°, 30°, 45° and 60°; see Figure 1). Critically, the letters of the fixated target word—as well as the rest of the sentence—were always in the canonical upright orientation. We employed 15° of rotation as a baseline rather than 0° to avoid a perceptual continuity between preview and target in the identity condition.² All the sentences were presented in UPPERCASE, the reasons being that: (a) in lowercase, some letters could be ambiguous when rotated (e.g., the rotated letter *b* can be confused with the letters *p* or *q*); (b) in lowercase words, low spatial frequency information spanning the whole word (e.g., compare bishop vs. BISHOP) would have been less disrupted for smaller than for larger rotation angles (i.e., the whole visual word form of lowercase words could be more informative for smaller than for larger rotation angles). Of note, the pattern of eye movements when reading sentences is similar when written in lowercase or uppercase (see Perea et al., 2017).

The predictions were the following. First, if, as hypothesized by the LCD model (Dehaene et al., 2005), letter detectors are hindered by rotations in the first stages of orthographic-lexical processing only when they are above 40° to 45°: (a) we would expect shorter fixation durations on the target word after an identity preview than after an unrelated preview, and to a similar degree, at 15° and 30° rotations (i.e., well below 45°); and (b) we would expect a dramatic drop in preview effects when the letters of the parafoveal preview are rotated 45° or 60°. Second, if, as hypothesized by the SERIOL model (Whitney, 2002), letter rotation gradually decreases the activation reached by the letter nodes, one

would expect progressively smaller advantages of the identity preview as a function of rotation angle. A final scenario is that the orthographic coding system is widely tuned to rapidly identify rotated and/or distorted letters with little cost (see Hannagan et al., 2012; Yang & Lupker, 2019; and Perea et al., 2018, 2020)—in this case, we would expect an advantage of the identity condition over the unrelated condition not only for the smaller rotation angles (15° and 30°) but also for the larger rotation angles (45° and 60°).

Method

Participants

Forty participants (18 female) were recruited at Aix-Marseille University. The participants were all native French speakers and received either course credit or monetary compensation (€10/h). All participants reported normal or corrected-to-normal vision and ranged in age from 18 to 38 years old ($M = 21.1$ y.o., $SD = 4.11$). They were naïve to the purpose of the experiment and signed an informed consent form before starting the experiment. Ethics approval was obtained from the Comité de Protection des Personnes SUD-EST IV (No. 17/051), and this research was carried out in accordance with the provisions of the World Medical Association Declaration of Helsinki.

Apparatus

Stimuli were displayed using OpenSesame (Mathôt et al., 2012) with each sentence occupying a single line. Eye movements were recorded with an EyeLink 1000 system (SR Research, Mississauga,

² Keep in mind that any difference between a 15° and the canonical condition could be attributable to the congruency of format (both upright) that would induce the 0° rotation angle.

ON, Canada) with high spatial resolution ($.01^\circ$) and a sampling rate of 1,000 Hz—this device only recorded the right eye movements. The sentences were presented on a 20-in. ViewSonic CRT monitor with a refresh rate of 150 Hz and a screen resolution of $1,024 \times 768$ pixels (30×40 cm). Stimuli were presented in upper case 18-point monospaced font (droid sans mono; the default monospaced font in OpenSesame) and the text was presented in black ($.15 \text{ cd/m}^2$) on a white background (21.70 cd/m^2); crosses and dots were presented with the same colors. Participants were seated 86 cm from the monitor, such that every three characters equaled approximately 1° of visual angle. The maximum validation error accepted was $.63^\circ$ of visual angle. A chin rest and forehead rest were used to minimize head movements.

Materials

We created 250 sentences in French, each contained an average of nine words (range 7–10). Sentences were created so as to minimize semantic predictability of the target words. This was verified via a cloze task in which the initial part of each sentence—until the word preceding the target word—was presented to four naïve individuals that were asked to predict the following word—none of these individuals took part in the experiment. The percentage of words that was predicted from the previous context was very low ($<1\%$). The pretarget and target words had an average length of seven letters and an average Zipf frequency of 5.71 and 6.13, respectively (van Heuven et al., 2014). Word frequencies were obtained using the film subtitle frequencies of the Lexique2 database (New et al., 2004). For each target word, we created eight parafoveal previews: identity condition with (a) 15° of rotation, (b) 30° of rotation, (c) 45° of rotation, (d) 60° of rotation, and unrelated condition with (e) 15° of rotation, (f) 30° of rotation, (g) 45° of rotation, and (h) 60° of rotation (see Figure 1). The parafoveal previews were counterbalanced across eight lists following a Latin square design. Each participant received 30 trials in each of the eight conditions and the sentences were presented in a different random order. The complete set of sentences, including the parafoveal previews, is presented in the Appendix.

Procedure

The experimental session took place in a dimly lit room. Participants were first informed about the experiment procedure. Their task was to read sentences for comprehension on a computer screen while their eye movements were registered. They were told that there would be comprehension questions after each sentence. The experiment procedure began with a five-point calibration phase in which participants had to look at individual dots on the screen. This was followed by 10 practice sentences to familiarize them with the procedure. Each trial had the following arrangement. First, a drift correction dot located 12 pixels (to the right of the left edge) was displayed. Second, the sentence was presented once the participant steadily looked at the drift correction dot—the starting point of the sentence was located just to the right of the dot. The display change occurred when the participant's eyes crossed an invisible boundary located between the target word and the word before (see Figure 1). Another invisible boundary was defined close to the end of the sentence (50 pixels before), such that the sentence disappeared when the eyes crossed that boundary

(that is, when participants read the last word). Finally, participants were shown a yes/no comprehension question after each trial that allowed us to verify that they had read for comprehension—participants were asked to press a corresponding key on a gamepad.

Data Analyses

The critical idea underlying Rayner's (1975) boundary technique is that the parafoveal preview allows some preprocessing of the target word. Specifically, if the reader extracts useful information from the rotated-identity parafoveal preview than for the rotated-unrelated parafoveal preview, the processing time on the target word would be reduced when directly fixated, thus resulting in shorter fixation durations (see Rayner et al., 2012). We examined three early eye fixation measures on the target word (that is, the critical region): (a) the duration of the initial fixation on the target word (FFD); (b) the duration of a fixation on the target word when it was the only fixation on that word in first pass reading (SFD); and (c) the sum of fixation durations before leaving the target word (GD).

For the inferential analyses, we analyzed the eye fixation data with linear mixed effects (LME) and Bayes Factors. We employed the *lme4* (Bates et al., 2015, 2019), *lmerTest* (Kuznetsova et al., 2016), *car* (Fox & Weisberg, 2019), and the *BayesFactor* .9.12-4.2 (Morey & Rouder, 2018) packages in R (R Core Team, 2020). The models included two fixed factors: (a) preview-target relationship and (b) preview rotation angle—each factor was zero-centered. Subjects and items were the random factors in the models. Variables were log-transformed to reduce highly positive skew of fixation durations and keep the normality assumption of LME models. Regarding LME analyses, the maximal random structure model that converged for the three dependent variables was (Dependent Variable) $\sim$ Relationship $\times$ Angle $+ (1 + \text{Relationship}|\text{Subject}) + (1 + \text{Relationship}|\text{Item})$. The *p* values for each main factor and the two-way interaction were calculated using Wald's chi-square tests. In case of a significant interaction, we computed the effect of preview-target relationship for each rotation angle with *emmeans* (Lenth et al., 2020).

Results

Participants were quite accurate in responding to the comprehension questions (mean accuracy was .91; $SD = .0023$). Trials where the display change was triggered either early or late were removed from the analyses (less than 1%). The raw data were preprocessed using EyeLink algorithms to detect saccades, fixations, and eye-blinks—the deletion of eye-blinks resulted in .5% of the data lost in the target region. Trials where the target word was skipped in first pass were excluded from the analysis (9.78%). Fixation durations of less than 80 ms and more than 800 ms and gaze durations longer than 2,000 ms on the target word were removed from the analyses. Furthermore, we removed any fixation duration measures that were more than 2.5 standard deviations for measure, condition, and subject because they are unlikely to represent normal reading—we removed the 2.07% of FFD observations, the 3.23% of GD observations and the 1.95% of SFD observations. The average eye fixation measures per condition across participants are presented in Table 1.

Table 1*Means (in ms) and Standard Errors (in Parentheses) as a Function of Rotation Angle and Preview-Target Relationship*

Parafoveal preview							
Rotation angle	Preview-target relationship	First fixation duration		Single fixation duration		Gaze duration	
15°	Identity	240 (2)	21	243 (3)	26	303 (5)	19
	Unrelated	262 (2)		269 (3)		322 (4)	
30°	Identity	252 (2)	8	258 (2)	13	302 (4)	12
	Unrelated	260 (2)		271 (3)		314 (4)	
45°	Identity	254 (2)	5	261 (3)	7	307 (4)	10
	Unrelated	259 (2)		268 (3)		317 (4)	
60°	Identity	256 (2)	4	265 (3)	3	306 (4)	9
	Unrelated	260 (2)		268 (3)		315 (4)	

Note. Values in bold represent the significant parafoveal preview effects.

First Fixation Duration

On average, first-fixation durations on the target word were 9 ms shorter in the identity than in the unrelated condition, $\chi^2(1) = 25.616$, $p < .001$. We also found a main effect of angle, $\chi^2(3) = 14.840$, $p < .001$. More important, the interaction between preview-target relationship and rotation angle was significant, $\chi^2(3) = 34.392$, $p < .001$. As can be seen in Figure 2, this interaction reflected that first fixation durations on the target word increased as a function of rotation in the identity condition, as reflected by significant linear ($t = 6.029$, $p < .001$) and quadratic ($t = -3.02$, $p = .002$) contrasts. For unrelated pairs, the duration of first fixations did not change as a function of rotation (none of the contrasts approached significance, all $|ts| < 1.25$, $ps > .21$). The interplay between the two factors (i.e., preview-target relationship and angle) resulted in shorter first-fixation durations on the target word in the identity condition than in the unrelated condition when the letters of the preview were rotated 15° (21 ms; $b = -.089$, $SE = .012$, $t = -7.627$, $p < .001$) and, to a lesser degree, when rotated 30° (8 ms; $b = -.028$, $SE = .012$, $t = -2.418$, $p = .016$). Finally, the advantage of the identity condition did not reach significance when the letters of the preview were rotated 45° (5 ms) or 60° (4 ms; both $ts < -1.146$, $ps > .253$). To examine whether the null effects for the 45° and 60° were attributable to evidence of absence rather than absence of evidence, we computed Bayes Factors, as these offer valuable information to express our degree of certainty that the rotation angle did not affect the effect of preview-target relationship. For computing the Bayes Factors, we employed the *lmBF* function from the *BayesFactor* package with the default Cauchy distribution (centered around 0 and with a width parameter $\delta = .707$; see Rouder et al., 2009; Wagenmakers et al., 2017; for discussion). The Bayes Factors for the 45° and 60° rotations were $BF_{01} > 7.508$, thus providing strong evidence toward the null hypothesis.

Single Fixation Duration

The proportion of first-pass trials with a single fixation in the target word was .66. Single fixation durations were, on average, 12 ms shorter in the identity than in the unrelated condition, $\chi^2(1) = 43.875$, $p < .001$. The effect of angle was also significant, $\chi^2(3) = 30.230$, $p < .001$, as was the interaction between the two factors, $\chi^2(3) = 41.251$, $p < .001$. This interaction showed that single fixation durations on the target increased nonlinearly with the

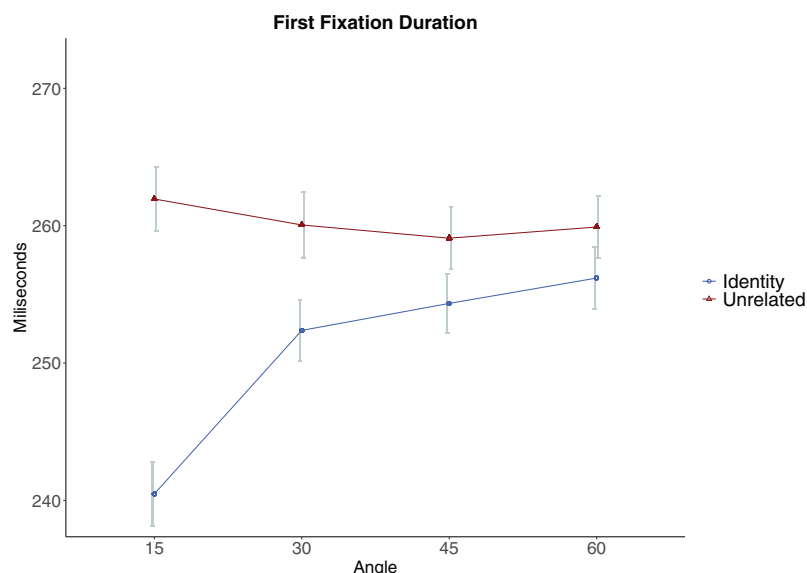
eccentricity of letter rotation in the identity condition—this was reflected by significant linear ($t = 8.414$, $p < .001$) and quadratic ($t = -3.205$, $p = .001$) contrasts (see Figure 3). The duration of single fixations for unrelated preview-target pairs remained constant independently of the rotation angle (none of the contrasts approached significance, all $|ts| < .76$, $ps > .45$). As a result of this interaction between preview-target relationship and angle, we found that single fixation durations were shorter in the identity condition than in the unrelated condition when the preview was rotated 15° (26 ms; $b = -.111$, $SE = .012$, $t = -9.357$, $p < .001$) and, to a lesser extent, 30° (13 ms; $b = -.045$, $SE = .012$, $t = -3.831$, $p < .001$). For the 45° rotation angle, the 7-ms advantage in the identity condition, approached significance ($b = -.014$, $SE = .012$, $t = -1.187$, $p = .079$), but the evidence favored, if anything, the null hypothesis, $BF_{01} = 1.742$. Finally, for the 60° rotation angle, we found strong evidence toward the null hypothesis (3-ms advantage in the identity condition; $b = -.014$, $SE = .012$, $t = -1.187$, $p = .236$; $BF_{01} = 6.295$).

Gaze Duration

On average, gaze durations on the target word were 12 ms shorter in the identity than in the unrelated condition, $\chi^2(1) = 41.022$, $p < .001$. The effect of angle did not approach significance, $\chi^2(3) = 5.339$, $p = .149$, but, more importantly, the interaction between preview-target relationship and angle was significant, $\chi^2(3) = 20.258$, $p < .001$. This interaction showed that gaze durations on the target word increased linearly with the eccentricity of letter rotation in the identity condition—this was reflected by significant linear contrasts ($t = 8.414$, $p < .001$). The duration of gaze durations for unrelated preview-target pairs remained constant independently of the rotation angle (none of the contrasts approached significance, all $|ts| < 1.52$, $ps > .12$). As a result of this interaction between preview-target relationship and angle, we found a 19-ms advantage of the identity condition 15° rotation ($b = -.093$, $SE = .013$, $t = -6.985$, $p < .001$), which was reduced to 12 ms at the 30° rotation ($b = -.052$, $SE = .013$, $t = -3.863$, $p < .001$). For the 45° rotation, we found a small 10-ms advantage of the identity condition ($b = -.027$, $SE = .014$, $t = -2.000$, $p = .046$), which actually slightly favored the null hypothesis, $BF_{01} = 1.688$, whereas for the 60° rotation, we found strong evidence for the null hypothesis (9 ms; $b = -.019$, $SE = .014$, $t = -1.398$, $p = .163$; $BF_{01} = 6.134$; see Figure 4).

Figure 2

Mean First Fixation Durations by Condition (Identity Versus Unrelated) and Angle (15°, 30°, 45°, and 60°)



Note. See the online article for the color version of this figure.

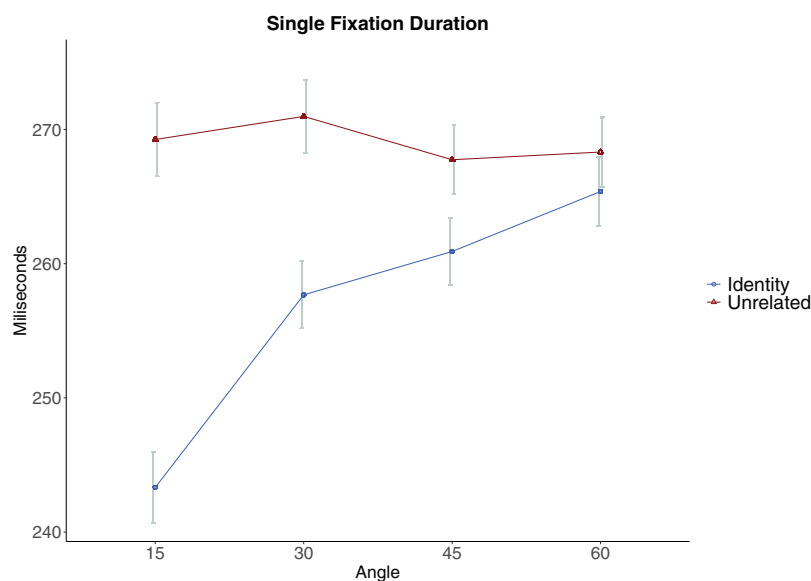
Discussion

The present experiment was designed to test whether letter detectors are severely hampered by rotations above 40° to 45° during sentence reading, as postulated by the LCD model (Dehaene et al., 2005), or whether this cost would emerge gradually as a function of increasing letter rotation, as hypothesized by the SERIOL

model (Whitney, 2002). To that end, we conducted a reading experiment with sentences in uppercase using Rayner's (1975) gaze-contingent boundary change paradigm in which we examined whether identity parafoveal previews were more effective than unrelated parafoveal previews when the previews were rotated at 15°, 30°, 45°, or 60° (see Figure 1 for illustration). We found a remarkably similar pattern of findings for all three eye fixation

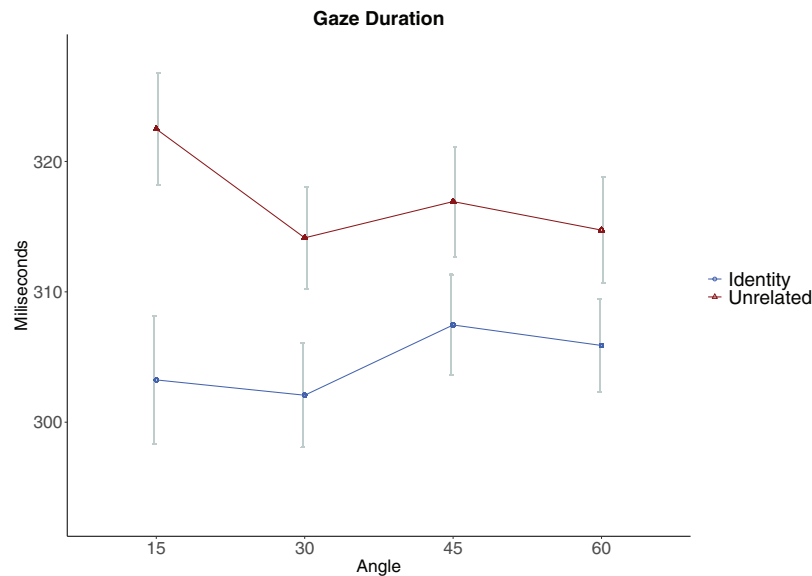
Figure 3

Mean Single Fixation Durations by Condition (Identity Versus Unrelated) and Angle (15°, 30°, 45°, and 60°)



Note. See the online article for the color version of this figure.

Figure 4
Mean Gaze Durations by Condition (Identity Versus Unrelated) and Angle (15°, 30°, 45°, and 60°)



Note. See the online article for the color version of this figure.

measures on the target word (i.e., first-fixation duration, single fixation duration, and gaze duration). Fixation durations were shorter when preceded by an identity preview than when preceded by an unrelated preview and, critically, this difference was modulated by rotation angle: the advantage of the identity preview condition was greatest when the letters of the preview were rotated 15° and it was numerically smaller but still robust at 30° (see Table 1). Instead, for the 45° rotations the advantage of the identity preview condition was only marginal, and the Bayes Factor analysis did not produce clear evidence in either direction—in particular for single-fixation durations and gaze durations. Finally, for the 60° rotations, there were no signs of an advantage of the identity preview condition—indeed, the evidence in favor of the null effect was strong (all $BF_{01} > 6.134$). Thus, the present results extend and qualify the findings obtained by Blythe et al. (2019) with a technique that specifically focused on early word recognition processes during sentence reading.

This pattern of results poses some problems to the predictions of the LCD model concerning the impact of letter rotation on word recognition (Dehaene et al., 2005). Although letter detectors appear to be severely disrupted at rotation angles of 45° or larger in the very early moments of processing, letter rotation also has an effect for rotation angles well below 45° in the parafovea. In the three eye fixation measures, we found a greater identity preview advantage for 15° than for 30° rotated previews. These findings, together with the data from Blythe et al. (2019) with 30° rotations, suggest that the disruption caused by letter rotations is not due to a sudden all-or-none shift at around the 40° to 45° boundary posited by the LCD model. Thus, the more parsimonious account of the current findings is that the amount of visual input reaching the letter detectors would decrease as a function of letter rotation, as proposed by the SERIOL model (see Whitney, 2002). This is also consistent with the increase in amplitude in ERP components that

reflect the initial contact with word-form features (e.g., N170) even for small rotations (22.5°) of individual letters in words (Kim & Straková, 2012). Second, although the evidence of an identity preview advantage for 45° rotations was marginal, we found strong evidence for a null effect for 60° rotations, thus again suggesting a gradual nature of the cost due to letter rotations on the letter detectors of the word recognition system. Taken together, this pattern would suggest that the identity preview advantage of rotated previews during sentence reading is structured by rotation angle (i.e., greatest at 15°, intermediate at 30°, marginal at 45°, and absent at 60°). Further experiments should be conducted to test whether this pattern of results also occurs with other manipulations (e.g., phonologically or graphemically preview-target related words).

Importantly, one might argue that our findings are at odds with the presence of robust masked priming effects with rotated words found by Perea et al. (2018) and Yang and Lupker (2019). However, there are some key methodological differences between the present and the just-mentioned studies. First, in the Perea et al. (2018) and Yang and Lupker (2019) experiments, the whole word was rotated, whereas in the present experiment we rotated the individual letters of the word (i.e., the manipulation was more extreme in the present experiment). Second, in above-cited masked priming experiments, (a) both primes and targets were presented foveally instead of parafoveal-then-foveal as in the present study (and the quality of visual information is higher in the fovea) and (b) both primes and the targets were presented rotated at the same angle. In the present experiment, the previews were presented parafoveally at different rotation angles, but the target word (as well as the rest of the sentence) was always presented in the canonical upright format. Thus, our participants were reading normally oriented sentences, thus avoiding the adoption of specific strategies that might arise when reading oriented text. To further elucidate these apparent

discrepancies, future research could examine the time course of masked priming words with individually rotated letters (e.g., using electrophysiological measures).

Our findings also have important implications for models of eye movement control in reading. At present, no current models of eye movement control during reading have explicit predictions on the processing of words composed of rotated letters (e.g., E-Z Reader, Reichle et al., 1998; SWIFT, Engbert et al., 2005; OB1-reader, Snell et al., 2018). This is because these models have focused on key elements of eye movement control (e.g., when and where to fixate the eyes) rather than on visual/sublexical/lexical processes. Our findings highlight the necessity of expanding these models to include early processes, including the mapping of visual information onto sublexical orthographic or phonological stages to offer a comprehensive picture of the reading process. As an example, the E-Z Reader model (Reichle et al., 1998, 2003) assumes that word identification is completed in two stages, reflecting early and late lexical processing. The first stage corresponds to the beginning of the identification of the orthographic form of the word, whereas full lexical access is reached in the second stage. The identity preview advantage observed in fixation durations when the letters of the rotated preview were 15° to 30° suggests that readers can complete the first stage of lexical processing (i.e., extract abstract letter codes) even when the individual letters of words are rotated 30° (i.e., readers are able to extract enough information to begin programming a saccade to the next word even when that information is extracted from a word with letters rotated 30°; see Angele et al., 2016, for a proposal of the mechanisms at play in parafoveal processing). As stated earlier, the cost of processing these rotated letters would be gradual (i.e., the larger the rotation angle, then smaller the information acquired in the parafovea).

In sum, skilled readers can efficiently convert parafoveal words composed of 15° to 30° rotated letters to a stable orthographic code during silent sentence reading, but not when the rotation angles are 45° or above. Thus, our data are in accordance with a gradual cost of letter rotations in the uptake of parafoveal information as a function of rotation angle, as posited by the SERIOL model of word recognition (Whitney, 2001, 2002). These findings should encourage the implementation of a more detailed account of the interface between visual form information and orthographic-lexical representations in computational models of eye movement control during sentence reading.

References

- Angele, B., Slattery, T. J., & Rayner, K. (2016). Two stages of parafoveal processing during reading: Evidence from a display change detection task. *Psychonomic Bulletin & Review*, 23(4), 1241–1249. <https://doi.org/10.3758/s13423-015-0995-0>
- Angele, B., Tran, R., & Rayner, K. (2013). Parafoveal–foveal overlap can facilitate ongoing word identification during reading: Evidence from eye movements. *Journal of Experimental Psychology: Human Perception and Performance*, 39(2), 526–538. <https://doi.org/10.1037/a0029492>
- Barnhart, A. S., & Goldinger, S. D. (2013). Rotation reveals the importance of configurational cues in handwritten word perception. *Psychonomic Bulletin & Review*, 20(6), 1319–1326. <https://doi.org/10.3758/s13423-013-0435-y>
- Bates, D., Mächler, M., Bolker, B., & Walker, S. (2019). lme4: Linear mixed-effects models using 'eigen' and S4 (1.1-20) [Computer software]. <https://CRAN.R-project.org/package=lme4>
- Bates, D., Mächler, M., Bolker, B., & Walker, S. (2015). Fitting linear mixed-effects models using lme4. *Journal of Statistical Software*, 67(1), 1–48. <https://doi.org/10.18637/jss.v067.i01>
- Blythe, H. I., Juhasz, B. J., Tbaily, L. W., Rayner, K., & Liversedge, S. P. (2019). Reading sentences of words with rotated letters: An eye movement study. *Quarterly Journal of Experimental Psychology: Human Experimental Psychology*, 72(7), 1790–1804. <https://doi.org/10.1177/1747021818810381>
- Clifton, C., Jr., Ferreira, F., Henderson, J. M., Inhoff, A. W., Liversedge, S. P., Reichle, E. D., & Schotter, E. R. (2016). Eye movements in reading and information processing: Keith Rayner's 40 year legacy. *Journal of Memory and Language*, 86, 1–19. <https://doi.org/10.1016/j.jml.2015.07.004>
- Cohen, L., Dehaene, S., Vinckier, F., Jobert, A., & Montavont, A. (2008). Reading normal and degraded words: Contribution of the dorsal and ventral visual pathways. *NeuroImage*, 40(1), 353–366. <https://doi.org/10.1016/j.neuroimage.2007.11.036>
- Dehaene, S., Cohen, L., Sigman, M., & Vinckier, F. (2005). The neural code for written words: A proposal. *Trends in Cognitive Sciences*, 9(7), 335–341. <https://doi.org/10.1016/j.tics.2005.05.004>
- Engbert, R., Nuthmann, A., Richter, E. M., & Kliegl, R. (2005). SWIFT: A dynamical model of saccade generation during reading. *Psychological Review*, 112(4), 777–813. <https://doi.org/10.1037/0033-295X.112.4.777>
- Fernández-López, M., Miraault, J., Grainger, J., & Perea, M. (2020, November 5). *How resilient is reading to letter rotations? A parafoveal preview investigation*. osf.io/h7uvj
- Fox, J., & Weisberg, S. (2019). An R companion to applied regression (R package version 3.0-8). <https://CRAN.R-project.org/package=car>
- Gomez, P., & Perea, M. (2014). Decomposing encoding and decisional components in visual-word recognition: A diffusion model analysis. *Quarterly Journal of Experimental Psychology: Human Experimental Psychology*, 67(12), 2455–2466. <https://doi.org/10.1080/17470218.2014.937447>
- Grainger, J., & Dufau, S. (2012). The front-end of visual word recognition. In J. S. Adelman (Ed.), *Visual word recognition* (Vol. 1, pp. 159–184). Psychology Press.
- Hannagan, T., Ktori, M., Chanceaux, M., & Grainger, J. (2012). Deciphering CAPTCHAs: What a turing test reveals about human cognition. *PLoS ONE*, 7(3), e32121. <https://doi.org/10.1371/journal.pone.0032121>
- Kim, A. E., & Straková, J. (2012). Concurrent effects of lexical status and letter-rotation during early stage visual word recognition: Evidence from ERPs. *Brain Research*, 1468, 52–62. <https://doi.org/10.1016/j.brainres.2012.04.008>
- Koriat, A., & Norman, J. (1984). What is rotated in mental rotation? *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 10(3), 421–434. <https://doi.org/10.1037/0278-7393.10.3.421>
- Kuznetsova, A., Brockhoff, P. B., & Christensen, R. H. B. (2016). lmerTest: Tests for random and fixed effects for linear mixed effect models (lmer objects of lme4 package; R package Version 2.0 –33). <http://CRAN.Rproject.org/package=lmerTest>
- Lenth, R., Singmann, H., Love, J., Buerkner, P., & Herve, M. (2020). emmeans: Estimated marginal means (R package version 1.4.8.). <https://CRAN.R-project.org/package=emmeans>
- Logothetis, N. K., & Pauls, J. (1995). Psychophysical and physiological evidence for viewer-centered object representations in the primate. *Cerebral Cortex*, 5(3), 270–288. <https://doi.org/10.1093/cercor/5.3.270>
- Mathôt, S., Schreij, D., & Theeuwes, J. (2012). OpenSesame: An open-source, graphical experiment builder for the social sciences. *Behavior Research Methods*, 44(2), 314–324. <https://doi.org/10.3758/s13428-011-0168-7>
- Morey, R. D., & Rouder, J. N. (2018). Package 'BayesFactor' (R Package version 0.9.12-4.2). <https://richarddmorey.github.io/BayesFactor/>

- New, B., Pallier, C., Brysbaert, M., & Ferrand, L. (2004). Lexique 2: A new French lexical database. *Behavior Research Methods, Instruments, & Computers*, 36(3), 516–524. <https://doi.org/10.3758/BF03195598>
- Perea, M., Marcet, A., & Fernández-López, M. (2018). Does letter rotation slow down orthographic processing in word recognition? *Psychonomic Bulletin & Review*, 25(6), 2295–2300. <https://doi.org/10.3758/s13423-017-1428-z>
- Perea, M., Rosa, E., & Marcet, A. (2017). Where is the locus of the lower-case advantage during sentence reading? *Acta Psychologica*, 177, 30–35. <https://doi.org/10.1016/j.actpsy.2017.04.007>
- Perea, M., Vergara-Martínez, M., & Gomez, P. (2015). Resolving the locus of cAsE aLtErNaTiOn effects in visual word recognition: Evidence from masked priming. *Cognition*, 142, 39–43. <https://doi.org/10.1016/j.cognition.2015.05.007>
- Perea, M., Vergara-Martínez, M., Marcet, A., Mallouh, R. A., & Fernández-López, M. (2020). When does rotation disrupt letter encoding? Testing the resilience of letter detectors in the initial moments of processing. *Memory & Cognition*, 48(5), 704–709. <https://doi.org/10.3758/s13421-020-01013-9>
- R Core Team. (2020). *R: A language and environment for statistical computing*. R Foundation for Statistical Computing. <http://www.R-project.org/>
- Rayner, K. (1975). Parafoveal identification during a fixation in reading. *Acta Psychologica*, 39(4), 271–281. [https://doi.org/10.1016/0001-6918\(75\)90011-6](https://doi.org/10.1016/0001-6918(75)90011-6)
- Rayner, K. (2009). The 35th Sir Frederick Bartlett Lecture: Eye movements and attention in reading, scene perception, and visual search. *Quarterly Journal of Experimental Psychology: Human Experimental Psychology*, 62(8), 1457–1506. <https://doi.org/10.1080/17470210902816461>
- Rayner, K., Inhoff, A. W., Morrison, R. E., Slowiczek, M. L., & Bertera, J. H. (1981). Masking of foveal and parafoveal vision during eye fixations in reading. *Journal of Experimental Psychology: Human Perception and Performance*, 7(1), 167–179. <https://doi.org/10.1037/0096-1523.7.1.167>
- Rayner, K., Pollatsek, A., Ashby, J., & Clifton, C. Jr. (2012). *Psychology of reading*. Psychology Press. <https://doi.org/10.4324/9780203155158>
- Rayner, K., Well, A. D., Pollatsek, A., & Bertera, J. H. (1982). The availability of useful information to the right of fixation in reading. *Perception & Psychophysics*, 31(6), 537–550. <https://doi.org/10.3758/BF03204186>
- Reichle, E. D., Pollatsek, A., Fisher, D. L., & Rayner, K. (1998). Toward a model of eye movement control in reading. *Psychological Review*, 105(1), 125–157. <https://doi.org/10.1037/0033-295x.105.1.125>
- Reichle, E. D., Rayner, K., & Pollatsek, A. (2003). The EZ Reader model of eye-movement control in reading: Comparisons to other models. *Behavioral and Brain Sciences*, 26(4), 445–526. <https://doi.org/10.1017/s0140525x03000104>
- Rouder, J. N., Speckman, P. L., Sun, D., Morey, R. D., & Iverson, G. (2009). Bayesian *t* tests for accepting and rejecting the null hypothesis. *Psychonomic Bulletin & Review*, 16(2), 225–237. <https://doi.org/10.3758/PBR.16.2.225>
- Slattery, T. J. (2016). Eye movements: From psycholinguistics to font design. In M. Dyson (Ed.), *Digital fonts and reading* (pp. 54–78). World Scientific. https://doi.org/10.1142/9789814759540_0004
- Snell, J., van Leipsig, S., Grainger, J., & Meeter, M. (2018). OB1-reader: A model of word recognition and eye movements in text reading. *Psychological Review*, 125(6), 969–984. <https://doi.org/10.1037/rev0000119>
- Tinker, M. A. (1958). Recent studies of eye movements in reading. *Psychological Bulletin*, 55(4), 215–231. <https://doi.org/10.1037/h0041228>
- Tinker, M. A., & Paterson, D. G. (1931). Studies of typographical factors influencing speed of reading. VII. Variations in color of print and background. *Journal of Applied Psychology*, 15(5), 471–479. <https://doi.org/10.1037/h0076001>
- van Heuven, W. J. B., Mandera, P., Keuleers, E., & Brysbaert, M. (2014). SUBTLEX-UK: A new and improved word frequency database for British English. *Quarterly Journal of Experimental Psychology: Human Experimental Psychology*, 67(6), 1176–1190. <https://doi.org/10.1080/17470218.2013.850521>
- Vasilev, M. R., Slattery, T. J., Kirkby, J. A., & Angele, B. (2018). What are the costs of degraded parafoveal previews during silent reading? *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 44(3), 371–386. <https://doi.org/10.1037/xlm0000433>
- Vinckier, F., Naccache, L., Papeix, C., Forget, J., Hahn-Barma, V., Dehaene, S., & Cohen, L. (2006). What” and “where” in word reading: Ventral coding of written words revealed by parietal atrophy. *Journal of Cognitive Neuroscience*, 18(12), 1998–2012. <https://doi.org/10.1162/jocn.2006.18.12.1998>
- Wagenmakers, E.-J., Marsman, M., Jamil, T., Ly, A., Verhagen, J., Love, J., Selker, R., Gronau, Q. F., Šmíra, M., Epskamp, S., Matzke, D., Rouder, J. N., & Morey, R. D. (2017). Bayesian inference for psychology. Part I: Theoretical advantages and practical ramifications. *Psychonomic Bulletin & Review*, 25(1), 35–57. <https://doi.org/10.3758/s13423-017-1343-3>
- Whitney, C. (2001). How the brain encodes the order of letters in a printed word: The SERIOL model and selective literature review. *Psychonomic Bulletin & Review*, 8(2), 221–243. <https://doi.org/10.3758/bf03196158>
- Whitney, C. (2002). An explanation of the length effect for rotated words. *Cognitive Systems Research*, 3(1), 113–119. [https://doi.org/10.1016/S1389-0417\(01\)00050-X](https://doi.org/10.1016/S1389-0417(01)00050-X)
- Yang, H., & Lupker, S. J. (2019). Does letter rotation decrease transposed letter priming effects? *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 45(12), 2309–2318. <https://doi.org/10.1037/xlm0000697>

(Appendix follows)

Appendix

Sentences Used in the Experiment

The parafoveal previews (identity preview; unrelated preview) are presented in parentheses. All words were presented in uppercase letters in the experiment.

Le petit cheval sauvage (sautera; animale) sautera par dessus la barrière.

Je pense que cette énorme (tomate; maison) tomate est bien mûre.

Tu vas avec ta gentille (copine; loupes) copine pour faire des courses.

Il avait vu un méchant (requin; dormir) requin dans la mer.

Elle marche avec de belles (dames; tenir) dames dans la grande avenue.

On aime regarder les beaux (bateaux; arbitre) bateaux dans le port.

Vous avez construit de grands (immeubles; commencer) immeubles dans la ville.

Elles nagent sur des grandes (vagues; rouges) vagues dans la mer.

Je vais acheter un nouveau (chapeau; malaise) chapeau pour ton mariage.

Tu as commandé un nouvel (appareil; claviers) appareil photo pour moi.

Elle vient de jeter ses nouvelles (vestes; souris) vestes par la fenêtre.

Ils iront voir leurs rares (amies; frire) amies dans la semaine.

Vous avez vraiment fières (allures; vendues) allures dans cette voiture.

Elles ont vu une grosse (pierre; vertes) pierre tomber de la falaise.

Nous avons vu ce surprenant (spectacle; enraciner) spectacle la semaine dernière.

Ta petite pièce grise (tombe; nuits) tombe de ta veste.

Tu voudras manger et aussi (dormir; partir) dormir avant de partir.

La coureuse va presque (arriver; devrait) arriver au sommet.

Je pense que tu devrais mieux (manger; cannes) manger pour ta santé.

Tu estimes avoir assez (dormi; chose) dormi pour cette nuit.

Elle va mieux après (avoir; serez) avoir raconté ses problèmes.

Vous avez beaucoup trop (fourni; nulles) fourni de preuves au juge.

Nous allons acheter trois (lampes; redire) lampes pour décorer.

Il me faudrait quatre (piles; ronde) piles pour le jouet.

Nous avons de tout petits (pots; airs) pots de fleur.

Vous venez de construire une grosse (caisse; soleil) caisse en bois.

Une oie a perdu ses jolis (oeufs; tours) oeufs dans le lac.

On aime votre joli (stylo; dunes) stylo tout neuf.

Il a vendu un étrange (tableau; pintade) tableau très cher.

Ne va pas dans cette étrange (chambre; acheter) chambre la nuit.

Ils sont sur la grande (falaise; louange) falaise de granite.

Elle lui jeta un terrible (regard; louves) regard dès son arrivée.

Vous avez toujours plusieurs (mouchoirs; abordions) mouchoirs dans votre poche.

Il navigue sur une large (barque; renard) barque sur la mer.

A ce compte autant (venir; comme) venir en pantalon.

Il faudra très vite (revenir; chutera) revenir voir ton docteur.

Il ne peut rien faire même (marcher; joueurs) marcher est douloureux.

Toi qui adores tant (boire; roses) boire de ce soda.

Je vais prendre quelques (parts; fours) parts de cette tarte.

Elle croit que tu aimes trop (couper; montre) couper de la viande.

La petite souris blanche (regarde; admirer) regarde le fromage.

Le petit enfant blond (saute; vitre) saute dans la rue.

Ton grand frère aime (acheter; secouer) acheter du pain frais.

Mon grand oncle est parti (vendre; poires) vendre un blouson.

Le jeune homme avait (perdu; creux) perdu sa fortune.

Le chevalier est encore (venu; toute) venu nous rendre visite.

Tu viens de partir (prendre; rideaux) prendre tes outils.

Le boulanger est prêt pour (malaxer; assiste) malaxer sa pâte à pain.

Vous êtes tous partis (danser; rasant) danser sur la piste.

Elle sait que tu aimes (rire; face) rire en ce moment.

Ils sont venus pour (rigoler; plieras) rigoler toute la soirée.

On sait que vous voulez surtout (chanter; adoucit) chanter cette chanson.

Nous savons que vous aimez beaucoup (jouer; germe) jouer du piano.

Cet enfant aime vraiment (lire; mais) lire des livres.

Ce petit monsieur va encore (grimper; alarmes) grimper ce matin.

La pilote va vite (freiner; adultes) freiner dans le virage.

Tu vas utiliser cette grande (planche; montres) planche de bois.

Elle nous offre sa belle (bague; tordu) bague de mariage.

Elle avait vu un gros (mouton; poteau) mouton sur une île.

Il vient de dire les dernières (phrases; argents) phrases de son discours.

Cette petite surface habitable (semblait; abattues) semblait bien décorée.

Leur grande croisière atypique (fait; irez) fait rêver leur voisin.

Votre paire de lunettes commandée (arrivera; orphelin) arrivera dans deux jours.

(Appendix continues)

Ce beau riz fermenté (donnera; lecture) donnera du bon saké.

Le joli petit pâton arrondi (formera; fragile) formera un beau pain.

Le sommeil du lion balafre (aura; fade) aura duré très longtemps.

Son ami très sceptique (chercha; galerie) chercha une autre solution.

La grande table calée (tient; tapis) tient bien droit.

Leur grande école laïque (semble; taille) semble vraiment réputée.

Le passage du pont enneigé (devra; sache) devra être interdit.

Un homme a préparé un infâme (dessert; percera) dessert dans la cuisine.

Ce travail au rythme soutenu (reste; boire) reste difficile à tenir.

Les gants de ski chauffant (sont; avez) sont bien agréables.

Un plan de maison simpliste (arrange; paniers) arrange tout le monde.

Une petite ère glaciaire (approche; barreaux) approche à grand pas.

Son point de vue imposé (agace; ongle) agace les élèves.

Son triste sentiment éprouvé (faisait; serions) faisait peine à voir.

La petite lame émoussée (taillera; paisible) taillera toujours aussi bien.

Ce dur règlement appliqué (donne; tuyau) donne de bons résultats.

Voir un mariage princier (fera; suis) fera rêver les gens.

Elle parla avec un influant (locuteur; gouverna) locuteur durant des heures.

Il a pris le mauvais (livre; pause) livre dans le tiroir.

Leur grand projet futuriste (aborde; combat) aborde de nouveaux problèmes.

Le cavalier amena son brave (cheval; enlace) cheval à la victoire.

Il découvrit ce troublant (morceau; affecta) morceau dans sa poche.

Les deux femmes prévoient un alléchant (accord; braver) accord ce soir.

Cette mignonne plante moisie (pousse; paquet) pousse malgré le froid.

Elle sentit une brusque (larme; garer) larme couler sur sa joue.

Ils ont acheté un soyeux (tissu; tarda) tissu chez le marchand.

Tant de rires joviaux (font; vais) font plaisir à entendre.

De gros nuages orageux (grondent; festival) grondent dans le ciel.

Cette petite personne loquace (parle; angle) parle vraiment bien.

Des grands hommes résignés (avancent; ampoules) avancent sous la pluie.

Quelques jeunes filles ennuyées (racontent; papillons) racontent des histoires.

Le pianiste ajoute une subtile (note; menu) note dans son récital.

Son long plan saugrenu (arriva; utiles) arriva à ses fins.

Des jeunes stagiaires investis (feront; allant) feront un bon travail.

Son grand air méprisant (cachait; panique) cachait une certaine timidité.

Je repensais au déchirant (souvenir; aggraver) souvenir de son départ.

Tous les vêtements blanchis (arrivent; fontaine) arrivent de leur machine.

Elle partage ce plausible (soucis; trucha) soucis avec ses parents.

Ce beau geste amical (engendre; baguette) engendre de belles émotions.

Certains gros sons brutaux (abiment; timides) abiment leur audition.

Plusieurs jeunes chiens têtus (courent; hommage) courent après le vélo.

La longue nappe tendue (forme; votez) forme une vraie table.

Le lourd sac chargé (alourdi; dispute) alourdi le coffre.

Le petit enfant honteux (rentre; unique) rentre sans un bruit.

Ses parents adoptèrent une sévère (attitude; enrichis) attitude après cette bêtise.

Ton célèbre roman policier (renvoie; poumons) renvoie à des faits réels.

Votre jolie carte postale (plait; sport) plait à toute ma famille.

Toutes les touches visibles (percutent; mannequin) percutent les cordes du piano.

Il vécut un insolite (moment; presse) moment avec son ami.

Il semblait jouer sur un banal (piano; nagea) piano durant le concert.

Elle a quitté son odieux (mari; taxe) mari pour être libre.

Ils habitent un hostile (lieu; bise) lieu au bout de la rue.

Il lui autorise un ultime (tournoi; daterai) tournoi dans la semaine.

Son superbe sourire éclatant (rayonne; baleine) rayonne de bonheur.

Ils ont partagé un intense (sourire; sources) sourire durant la danse.

Les mariés ont prévu trente (convives; coiffait) convives pour le repas.

Cet homme a vu un ignoble (cafard; casser) cafard sous son lit.

Son pauvre air navré (change; examen) change quand tu es présent.

Il a dormi dans cet abrupte (cachot; phrase) cachot chaque nuit.

Tu sens venir une intense (fatigue; emplois) fatigue dès maintenant.

Il regarda les cages et fortement (envoya; cabine) envoya le ballon.

Elle saisit une occasion et amicalement (invite; cahier) invite ses collègues.

Il parlait et consciemment (affichait; synergies) affichait ses opinions.

Elle réussit à lever ce lourd (marteau; gravure) marteau plusieurs fois.

Adrien est heureux et gaïement (chantonne; clochette) chantonne son air favori.

Ce livre très technique (fascina; globule) fascina ce jeune enfant.

Ils mirent en ordre cette large (faction; cracher) faction très armée.

Prenez en compte cette forte (amplitude; chalutier) amplitude pour atterrir.

Je veux visiter ces effroyables (caves; accru) caves du château.

Tu cuisines une fantastique (tarte; rhums) tarte à la rhubarbe.

Il faut faire assez (mariner; neigera) mariner pour le manger.

Il va falloir certainement (ajuster; ondoyer) ajuster ces boutons de manchette.

Les chevaliers voulaient bravement (prier; ceint) prier pour leur victoire.

Le juge de manière juste (accorda; sonneur) accorda la liberté au malheureux.

Nous devons soigner cette pathologique (toux; tact) toux avant Noël.

Je veux acheter ces charmants (tableaux; alluvion) tableaux dès demain.

Je transporte un fragile (carton; migrer) carton de décorations africaines.

Ton frère effectue de fières (actions; anoblir) actions dans son travail.

Il frémit peu puis (actionna; abattoir) actionna le levier.

Je ne peux objectivement (consentir; chapelles) consentir à cet accord.

Tu peux demander à son proche (cousin; infusa) cousin un avis.

Mon colocataire aime les franches (rigolades; immortels) rigolades avec ses amis.

Ta mère *SE* tient debout (face; lion) face au vent.

Parmi ces gens peu (acceptent; confiance) acceptent cette dure vérité.

Mon chien a environ (pris; rite) pris la moitié de sa gamelle.

Cet événement est historiquement (compris; visages) compris comme le commencement.

Il aurait voulu calmement (attendre; tringles) attendre la fin du récit.

Dors bien et demain (apportera; dictature) apportera de nouvelles perspectives.

Cette femme bien gentille (proposa; douille) proposa son aide.

Son grand père autrefois (rencontra; dominante) rencontra une bête velue.

Elle a obtenu une couverte (toiture; bouture) toiture pour cet hiver.

Je pense que demain (sera; ayez) sera un jour exceptionnel.

Je te voue une éternelle (affection; embauches) affection pour que tu restes.

Voilà encore une étrange (substance; attendons) substance à ne pas consommer.

Il faut procéder à un écologique (changement; tapisserie) changement pour la suite.

Elle a entendu une faible (voix; urne) voix murmurer à mon oreille.

Chaque nuit une étrange (vision; biches) vision *SE* présente à moi.

Tu as construit cette blanche (cabane; rester) cabane de tes mains.

Il adore son fraternel (cadeau; rivait) cadeau de Noël.

Il donna un mortifère (conseil; jubiler) conseil à son ami.

Elle effectue son premier (calibrage; robotique) calibrage sur cet appareil.

Ce sera le dernier (patron; coller) patron de ma vie.

Ce grand oiseau brièvement (admire; encode) admire un poisson dans le lac.

Toute cette pagaille énorme (engendrera; aérobiques) engendrera une fin terrible.

Il *SE* baigne dans la solaire (piscine; ligoter) piscine des voisins.

Il est en retard et finement (essaye; anneau) essaye discrètement de rentrer.

Ils offrent une moisie (couverture; carnivores) couverture pour la nuit.

Il faut entreprendre une militante (action; pagode) action pour marquer les esprits.

Elle présente sa fière (invention; musiciens) invention au concours.

Il nettoie cette lumineuse (vitre; clore) vitre pour bien voir.

Derrière la porte il a fixement (entendu; sourires) entendu la conversation.

Les américains larguent une atomique (bombe; savon) bombe sur le Japon.

Elle exhibe une tranchante (lame; posa) lame pour frimer.

Prudence avec cette fragile (barge; grise) barge pour touristes.

Le gouvernement promeut cette nationale (industrie; vagues) industrie pour notre économie.

Elle a comme une étrange (envie; devin) envie de partir loin.

Fais attention à la petite (marche; manger) marche sur le palier.

Il avait remarqué une étrange (lueur; nuage) lueur sur la montagne.

Ce petit animal bizarre (ressemble; entonnoir) ressemble à une tortue.

Le gros chien avait fortement (mordu; canne) mordu cette personne.

La hauteur de cette grande (tour; rive) tour est impressionnante.

Tous les enfants sages (auront; alliez) auront une surprise.

(Appendix continues)

Nous avons vécu une sacrée (histoire; explorer) histoire ce weekend.

Elle a endommagé le superbe (collier; trouver) collier de son anniversaire.

Il disait que cette nouvelle (image; chien) image était parfaite.

Nous étions fiers de notre récent (record; visage) record sur la course.

Il a adoré les délicieux (plats; niche) plats de ma grand mère.

Il suffit juste de bien (retenir; horizon) retenir les consignes.

Cette équipe recueille de nombreux (joueurs; exemple) joueurs très réputés mondialement.

Les personnes peuvent toujours (partir; copier) partir quand elles le souhaitent.

On pouvait trouver de grosses (erreurs; abandon) erreurs dans ses explications.

Vous avez acheté une splendide (voiture; activer) voiture de course.

Il ne faut jamais (courir; radote) courir au bord de la piscine.

Il aurait fallu seulement (quitter; enfants) quitter ses affaires.

Si seulement tu pouvais vraiment (analyser; panneaux) analyser la situation.

Tous ses amis ont encore (voulu; dises) voulu venir à la soirée.

Il suffisait de purement (organiser; libraires) organiser tes devoirs à faire.

Le tout est de bien (regarder; anecdote) regarder ce qui est présenté.

Nous avons dégusté un succulent (repas; vitre) repas cette soirée là.

La vieille voiture rouge (fonce; tenir) fonce à vive allure.

Le nouveau conseil restreint (envisage; annuelle) envisage une réforme complète.

De nombreux pays libres (veulent; vaudras) veulent rejoindre les nations unies.

Elle aime créer de nouvelles (recettes; postures) recettes pour chaque occasion.

Le petit ventilateur orange (souffle; segment) souffle vraiment trop fort.

La table en fer forgé (rend; doux) rend beaucoup service.

Le chien du voisin avait encore (grandi; treize) grandi dans la nuit.

Tu dois suivre le seul (protocole; retombait) protocole obligatoire pour ce travail.

Nous pourrions devenir de grandes (stars; herbe) stars de la musique.

Tu dois écouter ce vieil (homme; pensa) homme qui donne des conseils.

Cette belle cité ancienne (repose; citron) repose sous les fonds marins.

Il a trouvé cette petite (merveille; cependant) merveille dans son jardin.

Cette recette est un grand (secret; violon) secret chez les cuisiniers.

Tous mes cousins éloignés (donnent; petites) donnent des nouvelles.

Ces petites affaires secrètes (choquent; accident) choquent tout le monde.

Les inscriptions aux nouveaux (cours; fleur) cours *SE* font dès maintenant.

Un retard aussi infime (serait; dirais) serait certainement pardonné.

Elles ont envoyé un grand (nombre; ruelle) nombre de lettres.

Le tout est de vraiment (persuader; campement) persuader le grand public.

Elle nous a dit de mieux (scruter; session) scruter la falaise.

Ils devraient tous justement (suivre; parlas) suivre ton exemple.

Il a découvert une obscure (grotte; ruines) grotte dans cette région.

Les élections font beaucoup (parler; preuve) parler en ce moment.

Le gagnant de cet étrange (concours; boissons) concours est toujours inconnu.

Il a fait un délicieux (cookie; rendre) cookie pour toute la famille.

Nous avons construit une grande (maison; rapace) maison sur les hauteurs.

Ils ne veulent plus (entendre; culturel) entendre parler de cette histoire.

Les rumeurs sur cette ville mythique (restent; trieuse) restent encore peu répandues.

Cet ordinateur est une rapide (machine; textuel) machine comparée aux autres.

Cette grande avancée majeure (renforce; surmonta) renforce nos prédictions.

Avance prudemment sur cette grande (route; jalon) route pendant la nuit.

Toute la police locale (surveille; bravoures) surveille ses faits et gestes.

Il a acheté de fameux (fruits; persil) fruits au primeur du coin.

Le petit chiot adorable (remue; buste) remue la queue sans arrêt.

Nous remporterons cette dernière (bataille; souhaits) bataille contre nos ennemis.

Il est important de bien (situer; bronze) situer le contexte.

Cette vieille loi injuste (subsiste; purifies) subsiste dans notre pays.

Cette nouvelle rue étroite (devrait; prenais) devrait être plus sécurisée.

Il faut voir ce superbe (spectacle; attachera) spectacle la semaine prochaine.

Nous avons connu une horrible (aventure; chasseur) aventure la nuit dernière.

Tu dois imaginer un immense (espace; orange) espace plein de mystères.

(Appendix continues)

Nous voyons toujours ce vieux (chantier; trancher) chantier devant notre maison.

Des billets pour ce fabuleux (stade; poney) stade sont en vente.

Cette fameuse nouvelle excitante (contient; maillons) contient des faits mensongers.

Cette famille a de charmants (enfants; cherche) enfants très bien éduqués.

Il *SE* réveillera avec un sublime (paysage; montres) paysage autour de lui.

La nouvelle série hilarante (promet; votera) promet un grand succès.

Chaque année ces effarants (touristes; tacticien) touristes saccagent nos littoraux.

Received July 29, 2020

Revision received November 5, 2020

Accepted November 19, 2020 ■

UNITED STATES POSTAL SERVICE® Statement of Ownership, Management, and Circulation (All Periodicals Publications Except Requester Publications)

1. Publication Title: *Journal of Experimental Psychology: Learning, Memory, and Cognition*

2. Issue Frequency: **Monthly**

3. Issue Number: **12**

4. Filing Date: **November 2021**

5. Annual Subscription Price: **Motor \$205 Indiv \$620 Inst \$2,394**

6. Complete Mailing Address of Known Office of Publication (Not printer) (Street, city, county, state, and ZIP+4®):
750 First Street NE, Washington, D.C. 20002-4242

7. Complete Mailing Address of Headquarters or General Business Office of Publisher (Not printer):
750 First Street NE, Washington, D.C. 20002-4242

8. Full Name and Complete Mailing Address of Publisher, Editor, and Managing Editor (Do not leave blank):
American Psychological Association, 750 First Street, NE, Washington, DC 20002-4242
Editor (Name and complete mailing address):
Aaron S. Benjamin, University of Illinois at Urbana-Champaign, 627 Psychology Bldg., 603 E. Daniel St., Champaign, IL 61820
Managing Editor (Name and complete mailing address):
Rose Sokol-Chang, American Psychological Association, 750 First Street, NE, Washington, DC 20002-4242

9. Owner (Do not leave blank. If the publication is owned by a corporation, give the name and address of the corporation immediately followed by the names and addresses of all individuals owning or holding 1 percent or more of the total amount of stock. If not owned by a corporation, give the names and addresses of the individual owners. If owned by a partnership or other unincorporated firm, give its name and address as well as those of each individual owner. If the publication is published by a nonprofit organization, give its name and address.)
Full Name: **American Psychological Association**
Complete Mailing Address: **750 First Street, NE
Washington, D.C. 20002-4242**

10. Tax Status (For completion by nonprofit organizations authorized to mail at nonprofit rates) (Check one):
☒ The purpose, function, and nonprofit status of this organization and the exempt status for federal income tax purposes:
☐ Has Not Changed During Preceding 12 Months
☐ Has Changed During Preceding 12 Months (Publisher must submit explanation of change with this statement)

PS Form 3526, July 2014 (Page 1 of 4) (See instructions page 4) PSN: 7530-01-000-9001 PRIVACY NOTICE: See our privacy policy on www.usps.com.

UNITED STATES POSTAL SERVICE® Statement of Ownership, Management, and Circulation (All Periodicals Publications Except Requester Publications)

16. Electronic Copy Circulation

	Average No. Copies Each Issue During Preceding 12 Months	No. Copies of Single Issue Published Nearest to Filing Date
a. Paid Electronic Copies	115	104
b. Total Paid Print Copies (Line 16a) + Paid Electronic Copies (Line 16a)	117	106
c. Total Print Distribution (Line 16b) + Paid Electronic Copies (Line 16a)	98.29%	98.11%
d. Percent Paid (Both Print & Electronic Copies) (16b divided by 16c x 100)		

☐ I certify that 95% of all my distributed copies (electronic and print) are paid above a nominal price.

17. Publication of Statement of Ownership
☒ If the publication is a general publication, publication of this statement is required. 1003 best printed
in the **December 2021** issue of this publication. ☐ Publication not required.

18. Signature and Title of Editor, Publisher, Business Manager, or Owner
Leon Hawkins
APA Circulation Manager
Date: **11/21/21**

I certify that all information furnished on this form is true and complete. I understand that anyone who furnishes false or misleading information on this form or who omits material or information requested on the form may be subject to criminal sanctions (including fines and imprisonment) and/or civil sanctions (including civil penalties).

13. Publication Title: *Journal of Experimental Psychology: Learning, Memory, and Cognition*

14. Issue Date for Circulation Data Below: **June 2021**

15. Extent and Nature of Circulation

	Average No. Copies Each Issue During Preceding 12 Months	No. Copies of Single Issue Published Nearest to Filing Date
a. Total Number of Copies (Net press run)	268	159
b. Paid Circulation (By Mail and Outside the Mail)	94	87
(1) Mailed Outside-County Paid Subscriptions Stated on PS Form 3541 (Include paid distribution above nominal rate, advertiser's proof copies, and exchange copies)		
(2) Mailed In-County Paid Subscriptions Stated on PS Form 3541 (Include paid distribution above nominal rate, advertiser's proof copies, and exchange copies)		
(3) Paid Distribution Outside the Mails Including Sales Through Dealers and Carriers, Street Vendors, Counter Sales, and Other Paid Distribution Outside USPS® (e.g., First-Class Mail®)	20	16
(4) Paid Distribution by Other Classes of Mail Through the USPS	1	1
c. Total Paid Distribution (Sum of 15b (1), (2), (3), and (4))	115	104
d. Free or Nominal Rate Circulation (By Mail and Outside the Mail)		
(1) Free or Nominal Rate Outside-County Copies Included on PS Form 3541		
(2) Free or Nominal Rate In-County Copies Included on PS Form 3541		
(3) Free or Nominal Rate Copies Mailed at Other Classes Through the USPS (e.g., First-Class Mail)		
(4) Free or Nominal Rate Distribution Outside the Mail (Carriers or other means)	2	2
e. Total Free or Nominal Rate Distribution (Sum of 15d (1), (2), (3), and (4))	2	2
f. Total Distribution (Sum of 15c and 15e)	117	106
g. Copies not Distributed (See Instructions to Publishers #4 (page #3))	151	53
h. Total (Sum of 15f and g)	268	159
i. Percent Paid (15c divided by 15f times 100)	98.29%	98.11%

*If you are claiming electronic copies, go to line 16 on page 3. If you are not claiming electronic copies, skip to line 17 on page 3.

PS Form 3526, July 2014 (Page 2 of 4)

Can Response Congruency Effects Be Obtained in Masked Priming Lexical Decision?

María Fernández-López, Ana Marcet, and Manuel Perea
Universitat de València

In past decades, researchers have conducted a myriad of masked priming lexical decision experiments aimed at unveiling the early processes underlying lexical access. A relatively overlooked question is whether a masked unrelated wordlike/unwordlike prime influences the processing of the target stimuli. If participants apply to the primes the same instructions as to the targets, one would predict a response congruency effect (e.g., *book-TRUE* faster than *fiok-TRUE*). Critically, the Bayesian Reader model predicts that there should be no effects of response congruency in masked priming lexical decision, whereas interactive-activation models offer more flexible predictions. We conducted 3 masked priming lexical decision experiments with 4 unrelated priming conditions differing in lexical status and wordlikeness (high-frequency word, low-frequency word, orthographically legal pseudoword, consonant string). Experiment 1 used wordlike nonwords as foils, Experiment 2 used illegal nonwords as foils, and Experiment 3 used orthographically legal hermit nonwords as foils. When the foils were orthographically legal (Experiments 1 and 3; i.e., a standard lexical decision scenario), lexical decision responses were not affected by the lexical status or wordlikeness of the unrelated primes, as predicted by the Bayesian Reader model and the selective inhibition hypothesis in interactive-activation models. When the foils were illegal (Experiment 2), consonant-string primes produced the slowest responses for word targets and the fastest responses for nonword targets. The Bayesian Reader model can capture this pattern, assuming that participants in Experiment 2 were making an orthographic legality decision (i.e., anything legal must be a word) rather than a lexical decision.

Keywords: lexical decision, masked priming, word recognition

Supplemental materials: <http://dx.doi.org/10.1037/xlm0000666.supp>

Since its inception in 1984, Forster and Davis's masked priming technique (Forster & Davis, 1984) has become a fundamental tool for examining the initial moments of letter and word recognition in cognitive psychology (see Forster, 2013; Grainger, 2008) and cognitive neuroscience (e.g., Event Related Potentials [ERPs]: Grainger & Holcomb, 2009; Magnetoencephalography: Lehtonen, Monahan, & Poeppel, 2011; functional Magnetic Resonance Imaging [fMRI]: Dehaene et al., 2001). In the standard setup, a briefly presented lowercase prime is preceded by a pattern mask and followed by an uppercase target stimulus until the participant's response so that readers are not usually aware of the prime. This allows researchers to examine various types of prime-target relationships (e.g., identity priming: *steal-STEAL* vs. *horse-STEAL*; orthographic priming: *steam-STEAL* vs. *green-STEAL*; phonological priming: *steel-STEAL* vs. *shell-STEAL*; morphological priming: *stole-STEAL* vs. *brink-STEAL*; associative/semantic priming: *thief-STEAL* vs. *nurse-STEAL*; translation priming: *robar-STEAL* vs. *venir-STEAL*). Importantly, this technique minimizes the strategic

processes that may occur with visible, unmasked primes (for fMRI evidence, see Dehaene et al., 2001; for modeling evidence, see Gomez, Perea, & Ratcliff, 2013).

Although the masked priming paradigm has been employed with a number of laboratory word identification tasks (e.g., lexical decision, naming, semantic categorization, same-different), most masked priming experiments have used the lexical decision task. Indeed, researchers have compiled large masked priming data sets with this task (see Adelman et al., 2014, for a megastudy of masked priming lexical decision using 27 form-related priming conditions), and most contemporary models of visual word recognition offer quantitative predictions at the behavioral level (i.e., response times [RTs] and accuracy) of masked priming effects in the lexical decision task (e.g., dual-route cascaded model: Coltheart, Rastle, Perry, Langdon, & Ziegler, 2001; spatial coding model [SCM]: Davis, 2010; multiple read-out model [MROM]: Grainger & Jacobs, 1996; Bayesian Reader model [BR model]: Norris & Kinoshita, 2008).

In the present article, we focus on a relatively neglected phenomenon in masked priming lexical decision that may allow us to test the predictions of leading models of visual word recognition: the response congruency effect with unrelated primes. The response congruency effect in lexical decision can be defined as the difference in RTs between a congruent trial (i.e., when the prime elicits the same response as the target; e.g., *bright-WINDOW* or *geuzel-POLCIR*) and an incongruent trial (e.g., *geuzel-WINDOW* or *bright-POLCIR*). The idea is that participants may apply the

This article was published Online First October 25, 2018.

María Fernández-López, Ana Marcet, and Manuel Perea, Departamento de Metodología and ERI Lectura, Universitat de València.

Correspondence concerning this article should be addressed to Manuel Perea, Departamento de Metodología and ERI Lectura, Universitat de València, Avenida Blasco Ibáñez, 21, 46010-Valencia, Spain. E-mail: mperea@uv.es

task instructions (e.g., an implicit categorization of the “yes” vs. “no” decision) not just to the targets but also to the masked primes. As [Loth and Davis \(2010b\)](#) indicated, the examination of response congruency effects in masked priming lexical decision allows us to scrutinize “the state of the lexicon after the prime stimulus has been processed up to an early stage” (p. 174). Thus, this effect may allow us to elucidate the processes underlying masked priming lexical decision.

Response congruency effects have been reported in other masked priming tasks. For instance, [Naccache and Dehaene \(2001\)](#); see also [Dehaene et al., 1998](#)) found faster responses in a numerical categorization task (“Is the number higher than 5?”) when the masked prime is from the same category as the target (e.g., prime = 7; target = 9) than when it is not (e.g., prime = 2; target = 9). Importantly, [Dehaene et al. \(1998\)](#) found evidence of a differential effect on motor preparation for congruent and incongruent trials using the lateralized readiness potential as a marker, thus favoring the idea that participants make implicit decisions to the prime stimuli. Similarly, [Forster \(2004\)](#) found a response congruency effect in a masked semantic categorization task for narrow categories (e.g., *answer-WINDOW* faster than *collie-WINDOW* for the category “type of dog”). To explain this response congruency effect, Forster stated, “An implicit decision about the category membership of the prime is reached before an explicit decision about the target is reached” (p. 283). If this implicit decision also takes place in masked priming lexical decision, one would expect faster “yes” responses to words after an unrelated word prime than after an unrelated nonword prime—and conversely, one would expect faster “no” responses to nonwords after an unrelated nonword prime than after an unrelated word prime. Critically, a leading computational model of visual word recognition, namely, the BR model ([Norris, 2006](#); see also [Norris & Kinoshita, 2008](#)), makes a straight prediction: There should be no effects of response congruency in masked priming lexical decision. The predictions from interactive activation models are less straightforward, although the presence-absence of a response congruency effect may be informative in deciding whether there is homogeneous or selective inhibition at the lexical level.

The BR model ([Norris, 2006](#)) assumes that words are represented in an orthographic and lexical perceptual space. Lexical decisions are based on the comparison of the evidence that the visual input is produced—or not—by a word. Information that increases the odds that the stimulus is a word will decrease the odds that the stimulus is a nonword until a criterion is achieved and a response is triggered. To simulate masked priming lexical decision, the BR model assumes that “evidence from both the prime and targets are integrated in reaching a decision” ([Kinoshita & Norris, 2012](#), p. 3). The masked prime generates evidence that the target is in a particular area of the perceptual space, thus predicting faster lexical decision responses to orthographically related pairs (e.g., *trie-TRUE*) than to control pairs (*fiok-TRUE*). Critically, if prime and target do not share any orthographic features, the prime will not alter the evidence for the target in lexical decision—note that in [Norris and Kinoshita’s \(2008\)](#) account of masked priming lexical decision, the information from the prime is not sufficient to provide specific evidence that the target is a word or not. That is, neither the masked word prime *food* nor the masked pseudoword prime *fiok* would have an effect on the responses to the target word *TRUE* because they are far apart in perceptual space. As Norris and

Kinoshita claimed, “there should be no response-congruence effects in lexical decision” (p. 442). What we should also note, however, is that the BR model assumes that priming effects depend on the decision required by the task, thus not precluding an effect of response congruence in other tasks. For instance, the BR model can capture an effect of response congruence in masked priming semantic categorization tasks (e.g., “Is the word an animal name?”). In a semantic categorization task, “the evidence consists of distributed semantic features that are diagnostic of category membership” ([de Wit & Kinoshita, 2014](#), p. 1738). Specifically, semantic features like <builds nests> in the masked prime *eagle* would constitute evidence toward a “yes” decision, whereas semantic features like <made of metal> in the masked prime *stapler* would provide evidence toward a “no” decision, thus producing a response congruency effect.

Alternatively, the IA model ([McClelland & Rumelhart, 1981](#)) and its successors (multiple read-out model: [Grainger & Jacobs, 1996](#); dual route cascade model: [Coltheart, et al., 2001](#); spatial coding model: [Davis, 2010](#)) assume that all units at the word level receive activation from the letter and visual feature levels. When the activation level of a word unit exceeds its resting level—which is a function of word frequency—it sends inhibition to all other word units. In the IA model, a “yes” response in lexical decision occurs when the activation level of a word unit exceeds some asymptotic activation threshold (see [Jacobs & Grainger, 1992](#), for the first implementation of the IA model to lexical decision), whereas a “no” response would be produced when none of the word units exceeds threshold after a certain processing duration (see [Coltheart, Davelaar, Jonasson, & Besner, 1977](#)). The IA model can easily capture masked orthographic priming effects: The processing of a prime (e.g., the pseudoword *trie*) would partially activate similarly spelled word units (e.g., *true*, *tree*, and *trip*), thus predicting faster lexical decision times to *trie-TRUE* than to *fiok-TRUE* (see [Perea & Rosa, 2000](#), for simulations). Using the original setting of the IA model, the activated word units inhibit all other word units regardless of whether they are orthographically similar or not (i.e., homogeneous lateral inhibition). That is, upon presentation of the masked prime *food*, its corresponding word unit would send inhibitory signals to other word units (e.g., *TRUE*). This inhibition would occur to a lesser degree when the prime is not a word (e.g., the prime *fiok* would produce less inhibition on the word unit *TRUE* than the prime *food*). This specification of the IA model predicts slower “yes” lexical decision times for those target words preceded by an unrelated word prime than when preceded by an unrelated pseudoword prime (e.g., *food-TRUE* slower than *fiok-TRUE*). More generally, the model predicts that “yes” lexical decision times for those target words preceded by an unrelated prime should be a direct function of the lexical activation generated by the prime: high-frequency word prime (*food-TRUE*) > low-frequency word prime (*glop-TRUE*) > pseudoword prime (*fiok-TRUE*) > consonant-string prime (*ghdk-TRUE*; see [Bayliss, Davis, Brysbaert, Luyten, & Rastle, 2009](#), for simulations showing this exact pattern). Thus, the original IA model predicts a response congruency effect—or, rather, an effect of the wordlikeness of the unrelated primes—for “yes” responses.

Importantly, the predictions in the IA model are dependent on the nature of the parameter settings. [Davis and Lupker \(2006\)](#) proposed a selective inhibition mechanism in which the activated

word units would only inhibit those word units that are close in lexical space (i.e., similarly spelled word units; e.g., *tree*, *tire*, *tore*, *free* for the word unit *true* but not for unrelated word units like *food*). In this scenario, neither the presentation of the high-frequency word prime *food* nor the presentation of the consonant-string prime *ghdk* would modify the activation level of the target word *TRUE*. Therefore, the selective inhibition hypothesis in the IA model would predict a null effect of response congruency for “yes” responses (see Bayliss et al., 2009, for simulations).

Finally, the predictions on “no” responses in lexical decision in the IA model may depend on how these responses are performed. If the temporal deadline for “no” responses is adjusted after an unrelated prime, as suggested by Forster (1998), one would not expect any response congruency effects for nonwords. That is, the temporal deadline for the pseudoword *UFEZ* would be the same regardless of whether the unrelated prime is a word (e.g., *true*) or not (e.g., *tybe*).

The empirical evidence of an effect of response congruency—or, more generally, an effect of the wordlikeness of the unrelated prime—in masked priming lexical decision is scarce and nonconclusive (see Table 1 for a summary). Although most experiments have failed to obtain an effect of the unrelated primes (Bayliss et al., 2009; Loth & Davis, 2010a, Experiment 4; Norris & Kinoshita,

2008; Perea, Fernández, & Rosa, 1998; Perea, Gómez, & Fraga, 2010; Perea, Jiménez, & Gómez, 2014, Experiment 1; Sereno, 1991), several experiments did find an effect (e.g., Jacobs, Grainger, & Ferrand, 1995; Loth & Davis, 2010a, Experiments 1, 2, and 3; Perea et al., 2014, Experiment 2). There are two procedural differences between these two types of outcomes. The null findings occurred when the target stimuli were presented once and the foils were orthographically legal (i.e., the most usual scenario in word recognition experiments). In contrast, the experiments that reported an effect from the unrelated primes either repeated the target stimuli several times (Jacobs et al., 1995; Perea et al., 2014, Experiment 2) or the foils were orthographically illegal (e.g., *SYKDD*; Loth & Davis, 2010a). To reexamine this issue in depth, we focused on the response congruency effect when the target stimuli are presented only once (i.e., the most common setting in word recognition experiments), and we varied the wordlikeness of the nonword targets: orthographically legal nonwords matched in sublexical characteristics with the words in Experiment 1, orthographically illegal nonwords in Experiment 2, and orthographically legal nonwords with no close neighbors in Experiment 3.

The main goal of the present study was to examine in detail under which conditions the effect of response congruency could

Table 1
Summary of Masked Priming Lexical Decision Experiments With Different Types of Unrelated Primes

Research article	Target type	Mean RT (in ms) Prime type				
		HFW	H/LFW	LFW	PW	NW
Sereno (1991)	HFW	621			623	
	LFW			694	702	
	PW		689		678	
Perea, Fernández, & Rosa (1998)	HFW	676		681	679	
	PW	810		811	810	
Norris & Kinoshita (2008, Exp. 1)	HFW	561			561	
	LFW			634	621	
	PW		672		682	
Bayliss et al. (2009, Exp. 1)	HFW	597		600	602	615
	LFW	688		669	679	686
Bayliss et al. (2009, Exp. 2)	HFW	581		584		
	LFW	668		666		
Loth & Davis (2010a, Exp. 4)	MFW		636		638	
	PW		794		793	
Perea, Gómez, & Fraga (2010)	W		635		641	
	PW		697		695	
Perea, Jiménez, & Gómez (2014, Exp. 1)	HFW/LFW		678		682	
	PW		727		726	
Jacobs, Grainger, & Ferrand (1995, Exp. 2) (Items repeated several times)	HFW	493			513	^a
	PW	515			498	^a
Loth & Davis (2010a, Exp. 1)	HFW	475				498 ^a
	NW	520				497 ^a
Loth & Davis (2010a, Exp. 2)	HFW	531				539 ^a
	PW	586				576 ^a
Loth & Davis (2010a, Exp. 3)	HFW	588				601 ^a
	PW	713				721 ^a
Perea et al. (2014, Exp. 2) (Items repeated several times)	HFW	590			601	^a
	PW	668			665	

Note. Wordlikeness of nonwords targets from Loth and Davis (2010a): Experiment 4 > Experiment 3 > Experiment 2. RT = Reaction Time; W = word; HFW = high-frequency word; H/LFW = high/low-frequency word; LFW = low-frequency word; MF = medium frequency; NW = illegal nonword; PW = pseudoword; Exp. = Experiment.

^a Significant effect from unrelated primes.

occur in masked priming lexical decision. Unlike previous published experiments, which focused exclusively on very few types of orthographically legal primes or on a single comparison of orthographically legal versus orthographically illegal primes, we chose four types of unrelated primes that differed in wordlikeness (60 items per condition): (a) a high-frequency word (e.g., *línea-JAULA*, *coche-GAULO* [in English: *line-CAGE*, *car-GAULO*]), (b) a low-frequency word (e.g., *censo-JAULA*, *boina-GAULO* [*census-CAGE*, *beret-GAULO*]), (c) an orthographically legal pseudoword (e.g., *bunte-JAULA*, *vomon-GAULO* [*bunte-CAGE*, *vomon-GAULO*]), and (d) a consonant string (e.g., *pdtnr-JAULA*, *lbctr-GAULO* [*pdtnr-CAGE*, *lbctr-GAULO*]). Thus, there were two unrelated word priming conditions and two unrelated nonword priming conditions. This allowed us to examine not only the overall response congruency effect but also the effect of the frequency of the word primes (high-frequency vs. low-frequency) and the effect of legality of the nonword primes (orthographically legal pseudowords vs. consonant strings).

Importantly, we also examined whether the difficulty of the word–nonword decision could modulate the response congruency effect in masked priming lexical decision. The rationale here is that the processes underlying lexical decision move faster toward the “yes” or “no” boundaries (i.e., evidence for one response or the other is accumulated faster) when words and nonwords are easily discriminable (e.g., using orthographically illegal nonwords as foils; see Ratcliff, Gomez, & McKoon, 2004, for empirical and modeling evidence). This may maximize the chances to observe an effect because of an implicit decision on the primes. For instance, the high-frequency prime *world* might be implicitly categorized as “yes,” and this would speed up the responses to the target word *CHAIR* relative to the consonant-string prime *wtsbn*. However, this prime-induced congruency effect may be more difficult to detect when the nonword foils are orthographically legal (see Loth & Davis, 2010b, for a similar reasoning). (One might argue that the size of effects is usually diminished when RTs are faster, but masked priming effects tend to be similar in magnitude for fast and slow responses; see Gomez et al., 2013). In Experiment 1, the foils were orthographically legal nonwords matched in subsyllabic elements with the target words (e.g., *RERTA*). This is the typical scenario in word recognition experiments. In Experiment 2, the word–nonword discrimination was made substantially easier by having orthographically illegal nonwords (*DPTME*) as foils, and hence maximizing the chances of a prime-induced congruency effect on target stimuli from an implicit “yes” or “no” decision to the prime. To anticipate the results, we found no signs of an effect from the unrelated primes in Experiment 1, whereas in Experiment 2, we found an effect of the unrelated primes on target performance, but only for consonant-string primes. Thus, one could argue that this latter finding may have been driven by a legality judgment decision (“If the target is orthographically legal, say yes”) rather than by lexical decision. To examine whether the effects found in Experiment 2 occurred because of a change in the nature of the task (i.e., orthographic legality rather than lexical decision) or because the words and nonwords were easily discriminable, we designed a third experiment. In Experiment 3, the foils had no close neighbors (e.g., *GAZUR*), but they were always orthographically legal. Therefore, although the word–nonword discrimination in Experiment 3 was easier than in Experiment 1,

lexical decisions needed to be made in terms of lexical activation rather than orthographic legality.

According to the BR model (Norris & Kinoshita, 2008), one would expect no effects of the lexicality or frequency of the unrelated primes in masked priming lexical decision for word and nonword targets. Therefore, the presence of an effect from the unrelated primes would pose some problems for Norris and Kinoshita’s (2008) account of masked priming lexical decision—a similar null effect is predicted by Forster’s (2004) account of masked priming lexical decision, in which lexical decisions cannot be generated by masked primes. The predictions from the original IA model hinge on whether they assume selective inhibition or not; in the former case, the model predicts no effects from the different types of unrelated primes, whereas in the latter, the model predicts a direct relation between the lexical decision times and the wordlikeness of the unrelated primes (see Bayliss et al., 2009, for simulations). We defer a thorough description of more sophisticated mechanisms for lexical decision responses in other IA models (i.e., MROM and SCM).

Experiment 1

Method

Participants. Twenty-four students from the Universitat de València (València, Spain), all native speakers of Spanish with normal or corrected to normal vision and no history of reading disorders, voluntarily took part in the experiment. This study was approved by the Experimental Research Ethics Committee of the Universitat de València. Before starting the experiment, all participants signed an informed consent form.

Materials. We selected 240 five-letter target words from the Spanish subtitle database (Duchon, Perea, Sebastián-Gallés, Martí, & Carreiras, 2013). The mean frequency was 16 occurrences per million (range = 5–40)—this corresponded to an average Zipf frequency of 4.15 (range = 3.74–4.61; see van Heuven, Mandera, Keuleers, & Brysbaert, 2014, for the advantages of using Zipf frequencies when describing word information). The mean orthographic Levenshtein distance 20 (Yarkoni, Balota, & Yap, 2008) was 1.41 (range = 1.00–2.65). We also created a set of 240 nonwords to act as foils. These nonwords were created from the target words using Wuggy (Keuleers & Brysbaert, 2010). To act as unrelated primes, we selected 120 high-frequency words ($M = 201$ occurrences per million, range = 49–1304; mean Zipf frequency = 5.14, range = 4.69–6.11) and 120 low-frequency words ($M = 2$ occurrences per million, range = 0–5; mean Zipf frequency = 3.26, range = 2.17–3.70)—all of them of five letters. Each target stimulus, presented in uppercase, was preceded by a lowercase prime that could be (a) an unrelated high-frequency word, (b) an unrelated low-frequency word, (c) a pseudoword (obtained using Wuggy [50% from high-frequency words and 50% from low-frequency words]), or (d) a five-letter consonant string. The complete list of prime-target pairs is available in Appendix A. To rotate the four priming conditions across all target words, we created four counterbalanced lists following a Latin square design. To familiarize the participants with the task, the 480-trial experimental phase (i.e., 60 items per condition) was preceded by a short practice phase composed of 16 trials (eight word trials and eight nonword trials) with the same characteristics as the experimental trials—these practice trials were not included in the analyses.

Procedure. Each participant was tested alone in a quiet room. DMDX software (Forster & Forster, 2003) was used to display the sequence of stimuli and to register the timing and the accuracy of the responses. In each trial, a pattern mask (i.e., a series of five numbers) was displayed for 500 ms in the center of a computer screen. The mask was replaced by a lowercase prime stimulus for 50 ms, and then the prime was replaced by an uppercase target stimulus. The target stimulus was displayed on the screen until the participant responded or 2 s had passed. All stimuli were presented in black (Courier New 14-pt font) on a white background. Participants were instructed to decide whether the target stimulus was a Spanish word or not by pressing the “yes” or the “no” key. Both speed and precision were stressed in the instructions. Instructions did not mention the existence of any lowercase stimuli. Sixteen practice trials preceded the 480 experimental trials. Participants did not receive feedback on RTs or error rates during the experiment. The session lasted for approximately 18 to 22 min. The raw data of the experiments are provided as [online supplemental materials](#).

Results and Discussion

Error responses and extremely short RTs (less than 250 ms: four responses) were omitted from the latency analyses—note that the deadline to respond was set to 2 s. The mean RTs for the correct responses and the accuracy in each experimental condition are presented in [Table 2](#).

To examine the influence of unrelated primes on the processing of word–nonword targets, we employed linear mixed effects (LME) models in R (R Core Team, 2017) using the *lme4* (Bates, Mächler, Bolker, & Walker, 2015) and *lmerTest* packages (Kuznetsova, Brockhoff, & Christensen, 2016), with two fixed factors: Target Lexicality (word, nonword) and Prime Type (high frequency word, low frequency word, pseudoword, consonant string). For type of prime, we created three orthogonal contrasts following the three research questions: (1) the effect of prime lexicality (word prime [high frequency word, low frequency word] vs. nonword [pseudoword, consonant string]), (2) the effect of prime word-frequency (high frequency word–prime vs. low frequency word–prime), and (3) the effect of the wordlikeness of the nonword primes (pseudoword vs. consonant string). Because of the normality assumption required by LME analyses, the raw RTs were inverse-transformed ($-1,000/\text{reaction time [RT]}$). There were 10,957 observations in the latency analyses. We chose the maximal random effects model whose structure converged. (Using untransformed RTs in the LMEs or using ANOVAs produced exactly the same pattern of results as that reported here). For the accuracy analyses, the responses were coded as binary values (1 =

correct, 0 = incorrect), and we used the *glmer* function in the *lme4* package.

Latency analyses. RTs were faster for word targets than for nonword targets (640 vs. 741 ms, respectively; $t = 9.128$, $p < .001$). None of the planned contrasts approached significance: prime lexicality ($t = .004$, $p = .997$), word–prime frequency ($t = .454$, $p = .65$), and nonword–prime wordlikeness ($t = 1.291$, $p = .209$)—none of the interactions between target lexicality and prime type approached significance (all t s < 1.485 , p s $> .137$).

Accuracy analyses. Accuracy was higher for word targets than for nonword targets (93.8 vs. 96.5, respectively; $z = 6.949$, $p < .001$). Neither the effect of prime lexicality ($z = 1.778$, $p = .075$) nor the effect of prime word-frequency ($z = .634$, $p = .526$) was significant—the interactions with target lexicality were not significant either (all p s $> .524$). The effect of nonword–prime wordlikeness was significant ($z = 2.896$, $p = .003$): Participants were more accurate when the prime was a pseudoword than when the prime was a consonant string (95.6 vs. 94.0, respectively)—this effect occurred similarly for word and nonword targets, as can be deduced from the lack of interaction between the two factors ($z = .808$, $p = .419$).

RT distributions. To further examine whether the LME analyses missed some subtle effects (e.g., a facilitative effect in the leading edge of the RT distributions that could have been canceled by an inhibitory effect in the higher quantiles), we computed the .1, .2, .3, .4, .5, .6, .7, .8, and .9 quantiles for each participant and condition, and then calculated the values for each quantile across participants (i.e., vincentiled averages). As can be seen in the vincentile plots of the RT distributions shown in [Figure 1](#), all priming conditions behaved similarly, thus corroborating the LME analyses.

In sum, the present experiment, using orthographically legal nonwords as foils and a moderately large number of pairs in each condition (60 items per condition), showed that decision times to “yes” and “no” responses were not modulated by the wordlikeness of the unrelated primes, thus replicating previous experiments with wordlike nonword foils (Bayliss et al., 2009; Loth & Davis, 2010a, Experiment 4; Norris & Kinoshita, 2008; Perea et al., 1998, 2010, 2014, Experiment 1; Sereno, 1991). The only effect from the unrelated primes occurred in the accuracy data: Accuracy was higher when the target stimulus was preceded by a pseudoword prime than when preceded by a prime composed of random consonants—note, however, that the size of the effect (although significant) was very small (95.6 vs. 94.0, respectively).

The BR model and IA model—under the selective inhibition hypothesis—can easily capture this pattern of findings. The question now is whether we could find an effect of the unrelated primes in masked priming lexical decision in an extreme scenario in which the word–nonword discrimination is very easy. Experiment 2 was parallel to Experiment 1, except that we employed illegal nonwords as foils. As the rate of evidence accumulation for a “yes” or “no” response with illegal nonwords will be substantially higher than with wordlike nonwords, there is more room for a response congruency effect to appear. The idea is that participants are more likely to make an implicit decision based on the prime when the lexical status of prime and target is easily discriminable (e.g., the prime *gktbc* would produce evidence toward a “no” decision, whereas a prime like *house* would produce evidence toward a “yes” decision). Indeed, in a series of unpublished experiments, Loth and Davis (2010a) reported a response congruency effect in masked priming lexical decision when using

Table 2
Mean Correct Response Times (in ms) and Accuracy (in Parentheses) for Words and Pseudowords in Experiment 1

Target lexicality	Prime type			
	High-frequency word	Low-frequency word	Pseudoword	Consonant string
Word	635 (96.7)	641 (97.1)	633 (96.9)	637 (95.3)
Pseudoword	734 (94.0)	730 (94.2)	733 (94.2)	735 (92.7)

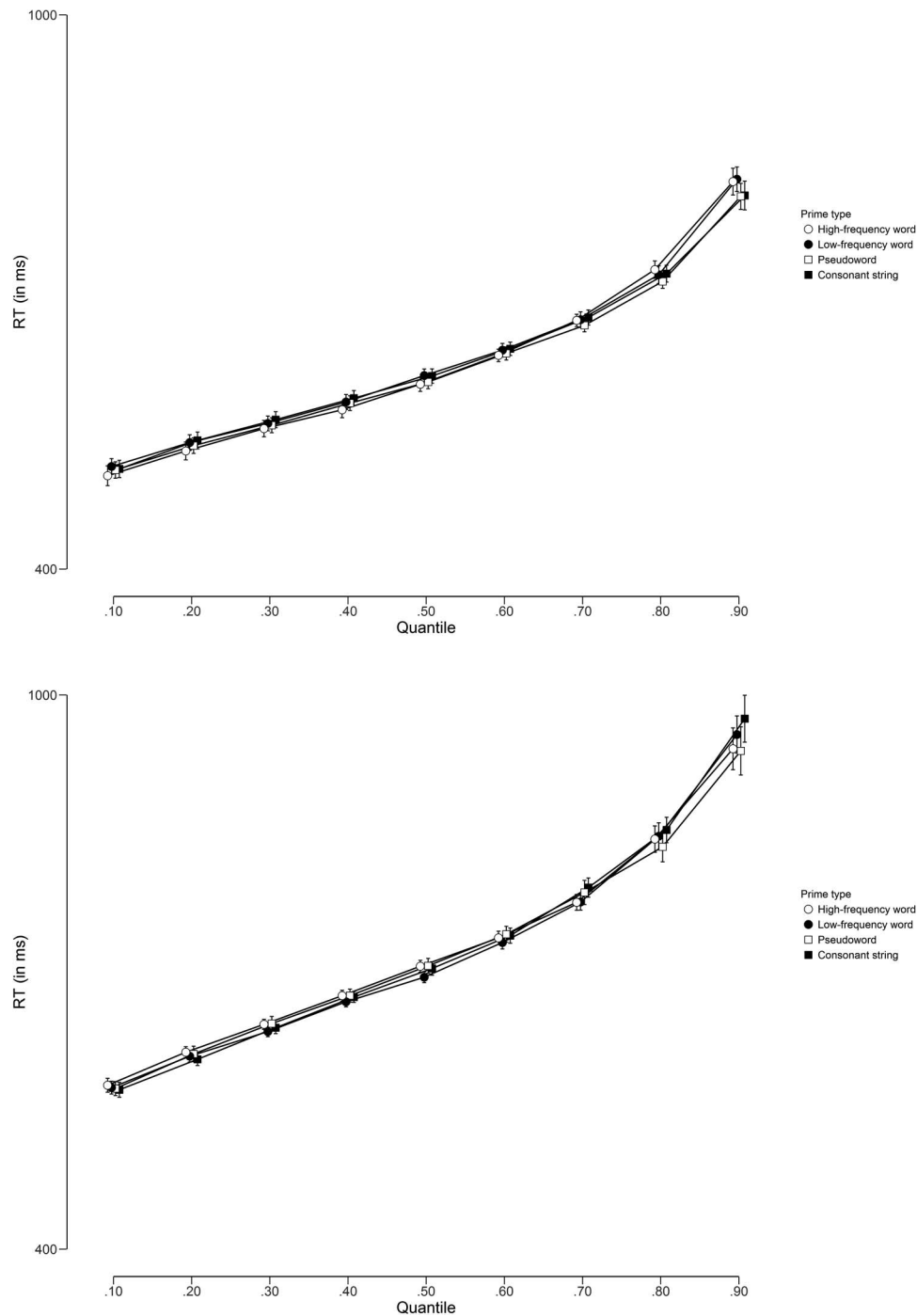


Figure 1. Group reaction time (RT) distributions for the four experimental conditions in Experiment 1 (upper panel: words; lower panel: nonwords) across quantiles (.1 to .9).

unwordlike nonword targets (e.g., faster “yes” responses for *order-CATCH* than for *dqrki-CATCH*; faster “no” responses for *dqrki-SYKDD* than for *order-SYKDD*). A limitation of Loth and Davis’s experiments, however, is that the two unrelated priming conditions (i.e., word primes vs. illegal primes) differed not only on lexical status (i.e., word vs. nonword) but also on their orthographic legality (legal vs. illegal).

Experiment 2

Method

Participants. Twenty-four new students from the same subject pool as in Experiment 1 participated in the experiment.

Materials. The set of stimuli was exactly the same as in Experiment 1 except for the nonword targets. Unlike Experiment

1, in which the foils were orthographically legal nonwords, in Experiment 2, all foils were illegal nonwords. All of the five-letter nonword foils contained an illegal trigram and a single vowel (e.g., *BEFRM*, *TCFRO*) so that they were hardly pronounceable—in a few cases (less than 9%), these nonwords had an orthographic neighbor (e.g., the nonword *chpja* has the very low-frequency word *chajá* [a Uruguayan dessert] as a neighbor). The complete list of nonword foils is available in [Appendix B](#).

Procedure. The procedure was the same as in Experiment 1.

Results and Discussion

As in Experiment 1, error responses and very short RTs (less than 250 ms: one response) were excluded from the latency analyses. The mean correct RTs and the accuracy in each condition are presented in [Table 3](#). The statistical analyses were parallel to those reported in Experiment 1. There were 11,215 observations in the RT analyses.

Latency analyses. RTs were faster for word targets than for nonword foils (546 vs. 576 ms, respectively; $t = 11.513$, $p < .001$).

Prime lexicality. Although the effect of prime lexicality was not significant ($t = .758$, $p = .449$), we found a significant interaction of prime lexicality and target lexicality ($t = 6.585$, $p < .001$): Lexical decision times on the target words were, on average, 10 ms faster when these were preceded by a word prime than when preceded by a nonword (pseudoword and consonant string) prime (541 vs. 551 ms, respectively; $t = 5.036$, $p < .001$), whereas lexical decision times on target nonwords were, on average, 7 ms faster when preceded by a nonword prime than when preceded by a word prime (573 vs. 580 ms, respectively; $t = 4.175$, $p < .001$).

Prime word-frequency. Neither the effect of the prime word-frequency ($t = .342$, $p = .732$) nor its interaction with target lexicality was significant ($t = 1.482$, $p = .138$).

Prime nonword-wordlikeness. The effect of prime nonword-wordlikeness was significant ($t = 3.573$, $p < .001$): RTs were faster when the prime was a pseudoword than when the prime was a consonant string (556 vs. 568 ms, respectively). This effect interacted with target lexicality ($t = 8.229$, $p < .001$): Lexical decision times on the target words were, on average, 28 ms faster when they were preceded by a pseudoword prime than when preceded by a consonant string (537 vs. 565 ms, respectively; $t = 7.704$, $p < .001$). In contrast, lexical decision times to target nonwords were, on average, 2 ms faster when the prime was a consonant string in comparison with pseudowords (572 vs. 574 ms, respectively; $t = 3.353$, $p < .001$). This latter difference, which was significant using inverse-transformed RT data, vanished when the analyses were conducted on the raw RT data ($t <$

1). As can be seen in the vincentile plot for nonword targets in [Figure 2](#), there was a substantial advantage of the consonant string condition over the pseudoword condition in the leading edge of the RT distribution that vanished in the higher quantiles—note that inverse transformations put more weight in the short than in the long RTs. To examine this issue in greater depth, we conducted a 9 (quantile) $\times 2$ (prime nonword-wordlikeness: pseudoword vs. consonant string) analysis of variance. The main effect of prime nonword-wordlikeness did not approach significance ($F < 1$), but critically, prime nonword-wordlikeness interacted with quantile, $F(8, 184) = 17.706$, $p < .001$: There was a substantial advantage of the consonant-string condition over the pseudoword condition in the leading edge of the RT distribution (29 ms in the .1 quantile, $p < .001$; 23 ms in the .2 quantile, $p < .001$; and 16 ms in the .3 quantile, $p = .009$), whereas the difference was (if anything) in the opposite direction at the highest quantile (–35 ms in the .9 quantile, $p = .108$)—we applied a Bonferroni correction in these simple effects tests. The nature of this interaction is discussed in the RT distributions section. (Note that the parallel analyses with word targets revealed a prime nonword-wordlikeness effect, $F[1, 23] = 59.83$, $p < .001$, that was approximately the same size across the quantiles [interaction, $F(8, 184) = 1.03$, $p = .41$]).

Accuracy analyses. We did not find any differences in accuracy between the words and nonwords ($z = .09$, $p = .926$).

Prime lexicality. The prime lexicality effect was significant ($z = 2.800$, $p = .005$): Participants were more accurate when the prime was a word than when the prime was a nonword (97.8 vs. 96.8, respectively). This effect interacted with target lexicality ($z = 2.450$, $p = .014$): For word targets, responses were more accurate when the unrelated prime was a word than when it was not a word (98.1 vs. 96.4; $z = 3.678$, $p < .001$), whereas this effect did not occur for nonword targets ($z = .236$, $p = .813$).

Prime word-frequency. Neither the effect of prime word-frequency ($z = .5$, $p = .619$) nor its interaction with target lexicality approached significance ($z = .34$, $p = .735$).

Prime nonword-wordlikeness. The effect of prime nonword-wordlikeness approximated significance ($z = 1.95$, $p = .052$). This effect interacted with target lexicality ($z = 2.660$, $p = .007$): For word targets, participants were more accurate when the prime was a pseudoword than when it was consonant string (97.6 vs. 95.1; $z = 3.482$, $p < .001$), but this difference did not approach significance for nonword targets ($z = .508$, $p = .611$).

RT distributions. As shown in the vincentile plot displayed in [Figure 2](#) (top panel), the RT distributions for the word targets were remarkably similar when preceded by an unrelated high-frequency word prime, an unrelated low-frequency word prime, or a pseudoword prime (i.e., the orthographically legal primes), whereas the RT distribution for the word targets preceded by a consonant-string prime showed a location shift to the right. Thus, the effect from consonant-string primes is not a simple response congruency effect (i.e., both consonant-strings primes and pseudoword primes presumably provide evidence for “no” responses in lexical decision, but the pseudoword primes behaved similarly to the word primes). Instead, what this pattern of data suggests is that participants took into account the orthographic characteristics of the primes regardless of their lexical status. This occurred fast enough so that an orthographically illegal prime produced evidence toward a “no” decision, whereas an orthographically legal prime produced

Table 3
Mean Correct Response Times (in ms) and Accuracy (in Parentheses) for Words and Nonwords in Experiment 2

Target lexicality	Prime type			
	High-frequency word	Low-frequency word	Pseudoword	Consonant string
Word	542 (98.2)	539 (97.9)	537 (97.5)	565 (95.1)
Nonword	577 (97.5)	583 (97.5)	574 (97.2)	572 (97.5)

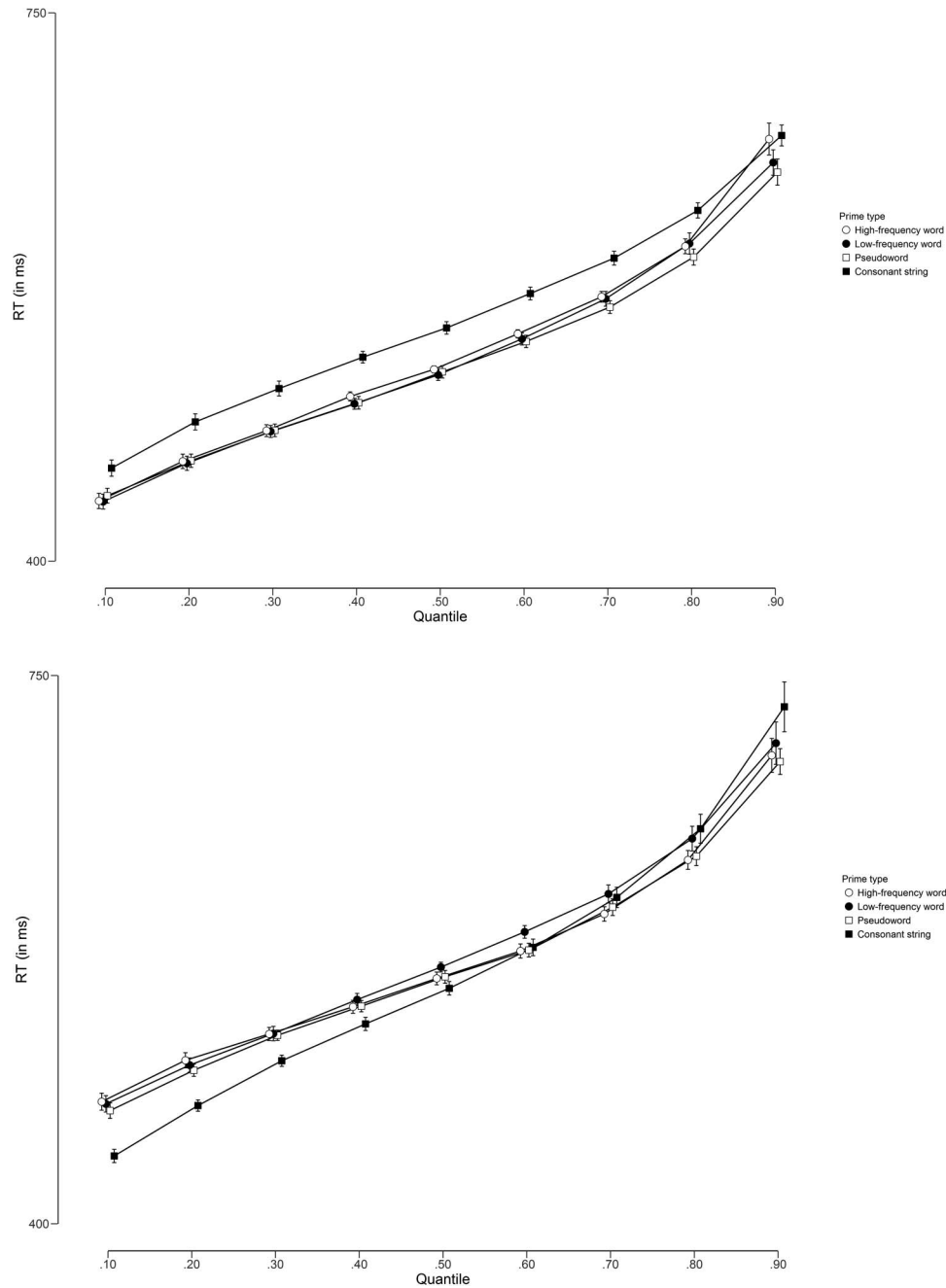


Figure 2. Group reaction time (RT) distributions for the four experimental conditions in Experiment 2 (upper panel: words; lower panel: nonwords) across quantiles (.1 to .9).

evidence toward a “yes” decision.¹ This reinforces the view that the participants in the Experiment 2 could have evaluated the legality of the stimuli rather than accessing the mental lexicon (i.e., a judgment legality task). Indeed, Forster, Mohan, and Hector (2003) showed that word-frequency effects (i.e., a marker of lexical access) are dramatically diminished in lexical decision experiments with illegal nonword foils—consistent with this view, the lexicality effect in the current experiment was dramatically smaller than in Experiment 1 (31 vs. 97 ms, respectively).

For the nonword targets, the RT distributions were again very similar when preceded by a high-frequency unrelated prime, a low-frequency unrelated prime, or a pseudoword prime (i.e., the orthographic legal primes). There is, however, a somewhat unusual pattern in the RT distribution of the consonant string condition: We found a large legality congruency effect in the initial quantiles

¹ We thank the reviewers and the editor for suggesting this explanation.

of the RT distribution (i.e., shorter “no” RTs to nonword targets when the prime was orthographically illegal than when the prime was orthographically legal) that vanished in the higher quantiles. This pattern is puzzling, as one would expect an effect to increase its size across quantiles (e.g., word-frequency effect; Ratcliff et al., 2004) or to be approximately the same size across quantiles (e.g., masked identity priming; Gomez et al., 2013). The existence of a large effect in the leading edge of the RT distribution that disappears in the higher quantiles has been documented in a numerical categorization task (“Is the digit higher than 5?”) with masked primes (e.g., prime = 7, target = 9 vs. prime = 3, target = 9; see Kinoshita & Hunt, 2008). Kinoshita and Hunt (2008) posited the existence of an “inhibitory control mechanism” that rejects an activated response when the cognitive system detects that “the (supraliminal) target is not the source of the response” (pp. 1331–1332). However, it is unclear how a prime can be discounted when it is not visible—the theories on prime discounting originated from studies using visible primes (e.g., see Huber & O’Reilly, 2003). Indeed, de Wit and Kinoshita (2014) acknowledged that “in masked priming, no prime discounting occurs because the prime is masked and hence the participants are unaware that the evidence accumulated from the prime comes from a difference source” (p. 1739). Clearly, the disappearance of the legality congruency effect in the higher quantiles for the nonword targets grants some explanation. The following is an admittedly ad hoc account that should be explored in more targeted studies.

Here, we propose a plausible mechanism within the structure of evidence accumulation modeling.² We term it the “probabilistic use of the prime” (PUP) hypothesis. The intuition is straightforward: Suppose that the observer does not use the information provided by the prime in every trial. Instead, the prime-target relationship is utilized in a proportion of trials (i.e., there is a probability greater than 0 and less than 1 that, in a given trial, the information provided by the prime is used). This hypothesis can be implemented in many different ways. However, given prior research on masked priming effects, we believe that assuming that prime-target congruency affects the Time of Encoding + Response (T_{er} parameter) is reasonable (see Gomez et al., 2013, for discussion). (Note that the diffusion model cannot disentangle the time of encoding from the time of response.) If the probability of using the prime-target relationship were 1 (PUP = 1), then the priming effect would produce a location shift in the RT distribution. Critically, as this probability goes down, the effect on the tail is progressively reduced (see Figure 3). Note that at relatively high PUPs (.85 in Panels C and D of Figure 3), there seems to be a location shift in the RT distributions. This matches the results with word targets in the current experiment and also the masked identity priming effects reported by Gomez et al. (2013; see also Perea, Marcet, Lozano, & Gomez, 2018). Importantly, when the PUP is low (.15 in Panels A and B of Figure 3), we can observe a pattern very similar to that found with nonword targets. That is, if prime-target legality congruency affects the T_{er} parameter in a relatively small proportion of trials (e.g., via an implicit decision to the prime stimulus), the model predicts an advantage of the congruent condition in the leading edge of the RT distribution that vanishes in higher quantiles. Why does the pattern of effects differ for word and nonword targets? We could speculate that the PUP might be related to the usefulness of the prime. If we assume that participants were basing their implicit decision on prime legality rather

than prime lexuality, three of the four priming conditions were congruent (i.e., high-frequency words, low-frequency words, orthographically legal pseudowords). Thus, in a large majority of trials, the orthographic legality of the primes would be helpful in reaching a “yes” decision (i.e., the implicit decision to the prime matches the decision to the target). If we assume that this affects mainly the T_{er} parameter of the model, one would obtain a legality congruency effect of similar size across the quantiles, as actually occurred. In contrast, for nonword targets, only one of four priming conditions was congruent. In this scenario, prime-target congruency would be substantially less helpful in reaching a “no” decision, and it may have been used in only a limited proportion of trials. Although admittedly ad hoc, the PUP hypothesis offers a reasonable account of the RT distributions for word and nonword targets in the present experiment (i.e., compare Figures 2 and 3). Importantly, the diminishing masked priming effect in the higher quantiles obtained by Kinoshita and Hunt (2008, Experiment 1) in a numerical categorization task can also be accommodated by the PUP hypothesis: In their experiment, only one of four priming conditions were congruent with the response.

Summary. In sum, the present experiment represents a demonstration of a response congruency effect in masked priming when the participants’ task requires the discrimination of words and illegal nonwords (e.g., *bunte-JAULA* [in English: *junt-CAGE*] produced faster “yes” responses than *pdtrv-JAULA* [*pdtr-CAGE*]; *ghdk-BEFRM* produced faster “no” responses than *fiok-BEFRM*). Importantly, these differences were completely absent with orthographically legal primes (i.e., high-frequency words, low-frequency words, and pseudowords behaved similarly; see Figure 2) and only occurred for consonant-string primes. This finding replicates and extends Loth and Davis’s (2010a) Experiment 1. In Loth and Davis’s Experiment 1, participants’ responses were 23 ms faster when the target word were preceded by a congruent prime (i.e., a high-frequency word; 475 ms) than when preceded by an incongruent prime (i.e., an illegal nonword; 498 ms)—we also found a 23-ms difference when comparing these two conditions.

What are the implications of this congruency effect for models of visual word recognition? Although computational models of visual word recognition have been used to simulate lexical decision experiments with illegal nonwords as foils (e.g., Loth & Davis, 2010b; Norris, 2009; Ratcliff et al., 2004), one might argue that under these circumstances, the participants’ task is not lexical decision per se, but rather an orthographic legality task (see Forster et al., 2003, for discussion)—note that we did not find a simple response congruency effect (i.e., unrelated pseudoword primes behaved just as unrelated word primes) but a legality congruency effect. To examine the involvement of lexical processing in Experiment 2, we calculated the Pearson correlation between the mean RT of each word and its Zipf frequency. The idea is that if participants are treating the task as an orthographic legality task, frequency effects should be very small. This analysis showed a negligible correlation between

² We are indebted to Pablo Gomez for suggesting this explanation. Although we use terms of the diffusion model account of lexical decision (Gomez et al., 2013; Ratcliff et al., 2004), Norris (2009) showed that this model can be subsumed within the BR model.

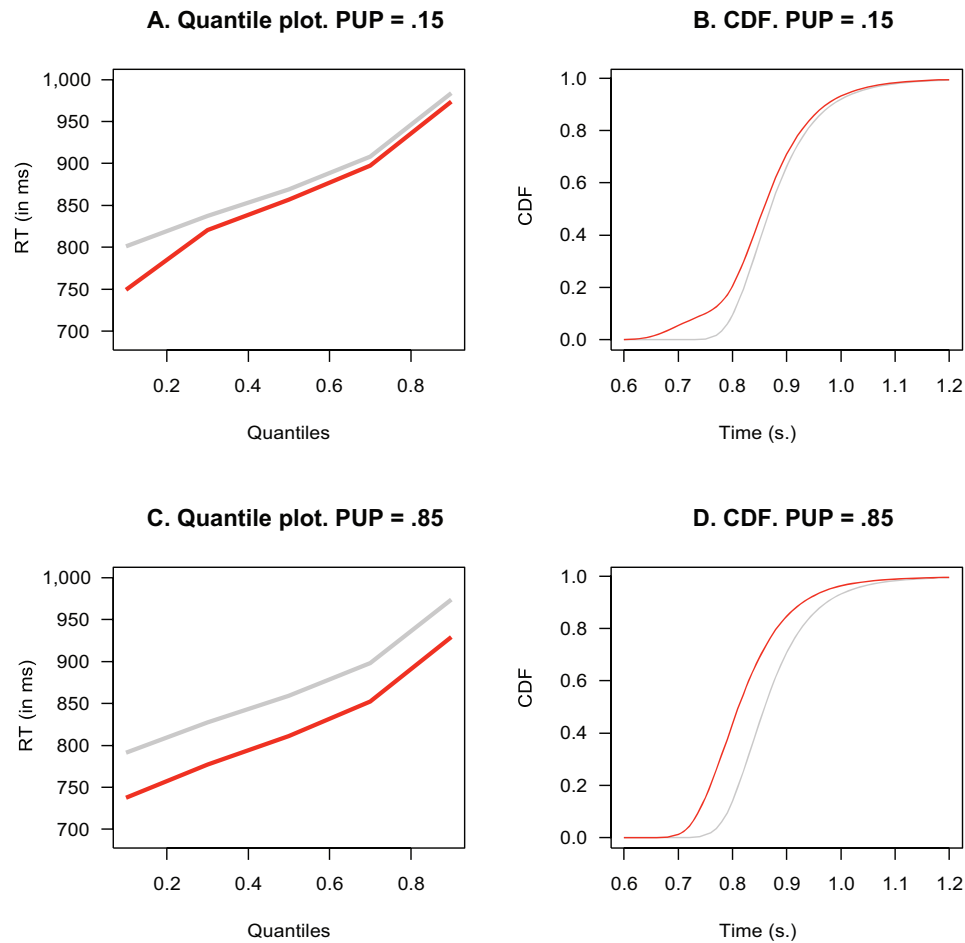


Figure 3. Quantile plots (left) and cumulative density functions (CDFs; right) for two conditions that differ in the probabilistic use of the prime (PUP): low proportion of trials (PUP = .15; top panels) versus high proportion of trials (PUP = .85, bottom panels). Prime-target relationship is assumed to affect the Time of Encoding + Response (T_{er} parameter) of the target. RT = Reaction Time. See the online article for the color version of this figure.

these two variables ($r = -0.04$, $p = .55$), thus supporting the view that the word–nonword task in the current experiment did not involve lexical access. (The parallel correlation coefficient in Experiment 1 was $r = -0.27$, $p < .001$.) Thus, we believe that it would be premature to make claims on whether the legality congruency effect in the current experiment poses serious problems for the masked priming account of the BR model or the selective inhibition hypothesis of the IA model in a standard lexical decision task—note that these two accounts would predict similar RTs for all priming conditions.

To create a more constraining scenario for the BR and IA models in an easy-to-perform lexical decision task without the risk of altering its lexical nature, we designed Experiment 3. Experiment 3 was similar to Experiment 2, in which the word–nonword discrimination was easy: The nonword foils did not have any close neighbors. Critically, all nonword foils were orthographically legal (GAZUR). In this scenario, lexical decisions need to be made in terms of lexical activation rather than on orthographic legality, so that the BR and the IA model—under the selective inhibition hypothesis—would predict a null effect of response congruency.

Experiment 3

Method

Participants. Twenty-four new students from the same subject pool as in Experiments 1 and 2 participated in the experiment.

Materials. The set of stimuli was exactly the same as in Experiments 1 and 2 except for the nonword targets. Unlike Experiment 2—in which foils were illegal nonwords—in Experiment 3, all foils were orthographically legal nonwords. We created the foils by changing two letters from Spanish words using Pseudo software (van Heuven, 2002). Then, we filtered the output by extracting only those nonwords with no close lexical neighbors (i.e., no one-letter substitution-letter neighbors, no transposed-letter neighbors, no addition-letter neighbors, and no deletion-letter neighbors). This produced a list of more than 300 orthographically legal hermit nonwords (e.g., *LEDUL*, *RUPUO*), and we randomly selected 240 nonwords from this set. The complete list of legal nonword foils is available in Appendix C.

Procedure. The procedure was the same as in Experiments 1 and 2.

Results and Discussion

As in Experiments 1 and 2, error responses and very short RTs (less than 250 ms: one response) were omitted from the latency analyses. The mean correct RTs and the accuracy in each experimental condition are presented in Table 4. The statistical analyses were parallel to those reported in the previous experiments. There were 11,167 observations in the RT analyses.

Latency analyses. RTs were, on average, 60 ms faster for word targets than for nonword targets (644 vs. 704 ms, respectively; $t = 5.775$, $p < .001$).

Prime lexicality. Neither the effect of prime lexicality nor its interaction with target type was significant (both t s < 1.26 , $ps > .21$). Indeed, lexical decision times were only 1 ms faster when the prime was a word than when the prime was a nonword (673 vs. 674 ms, respectively).

Prime word-frequency. Neither the effect of word-prime frequency nor its interaction with target type approached significant (both t s < 1 , $ps > .65$). Participants' responses were 1 ms faster when the target was preceded by a high-frequency word than when preceded by a low-frequency word (673 vs. 674 ms, respectively).

Prime nonword-wordlikeness. Neither the effect of prime nonword-wordlikeness ($t = 1.941$, $p = .065$) nor its interaction with target type was significant ($t = .403$, $p = .687$)—note that, for word targets, the difference between the pseudoword and consonant string conditions was only 4 ms (643 vs. 647 ms, respectively; $t = 1.275$, $p = .208$).

Accuracy analyses. Accuracy was higher for nonword targets than for word targets (97.7 vs. 96.2, respectively; $z = 4.44$, $p < .001$).

Prime lexicality. The effect of prime lexicality was significant ($z = 2.264$, $p = .02$): Participants were more accurate when the prime was a word than when the prime was a nonword (97.3 vs. 96.5, respectively)—this effect occurred similarly for word and nonword targets, as can be deduced from the lack of interaction between prime and target lexicality ($z = .789$, $p = .429$).

Prime word-frequency. Neither the effect of word prime frequency nor its interaction with target lexicality was significant (both z s < 1 , both $ps > .42$). Subjects were only 0.5% more accurate when the prime was a low-frequency word than when the prime was a high-frequency word (97.6 vs. 97.1, respectively).

Prime nonword-wordlikeness. Neither the effect of prime nonword-wordlikeness nor its interaction with target type approached significance (both z s < 1.24 , $ps > .21$). The difference in accuracy between pseudoword and consonant string primes is only of 0.5% (96.8 vs. 96.3, respectively).

RT distributions. As can be seen in the vincentile plots for both the word and nonword targets displayed in Figure 4, all priming conditions behaved similarly, thus corroborating the LME analyses.

Therefore, the present experiment showed that, when using orthographically legal hermit nonwords as foils, lexical decision times on the target stimuli were similar regardless of the characteristics of the unrelated primes. This pattern of data is similar to that found in Experiment 1, which employed orthographically legal nonwords matched with words in sublexical characteristics. Furthermore, the Pearson correlation between the mean RT of each word and its Zipf frequency in the current experiment was similar to that found in Experiment 1 ($r = -0.28$, $p < .001$ and $r = -0.27$, $p < .001$, respectively), thus showing that the two experiments involve lexical access. The only effect from unrelated primes occurred in the accuracy data: Accuracy was 0.8% higher when the target stimuli were preceded by a word prime than when preceded by a nonword prime (97.3 vs. 96.5, respectively).

The null finding in the latency data may be in contrast to Loth and Davis's (2010a) Experiment 2, which found that lexical decision times for word targets were faster when preceded by a high-frequency word prime than when preceded by an illegal nonword prime when the foils were pronounceable nonwords with no neighbors (e.g., *TEIR*, *DULEW*). However, the effect size in Loth and Davis's Experiment 2 was only 8 ms—the parallel difference in the current experiment was 3 ms (see Table 4). Thus, we prefer to remain cautious about this small effect.

General Discussion

The present experiments were designed to examine in detail whether the lexical status and wordlikeness of unrelated primes (high-frequency words, low-frequency words, orthographically legal pseudowords, and random consonant strings) could affect target processing in masked priming lexical decision. We did so in three different scenarios that varied in the difficulty of the word-nonword discrimination: orthographically legal nonwords matched in subsyllabic elements with the target words (Experiment 1), orthographically illegal nonword foils (Experiment 2), and orthographically legal nonwords with no neighbors (Experiment 3). When using orthographically legal nonwords, RTs to words and nonwords did not differ as a function of the wordlikeness of the prime stimuli (Experiments 1 and 3; see Figures 1 and 4), thus extending previous research with orthographically legal foils (see Table 1 for a summary) with a more thorough set of conditions. We only found an effect of the unrelated primes when the nonwords were orthographically illegal (Experiment 2) and restricted to the primes composed of consonant strings (see Figure 2). We now discuss the implications of these findings for leading models of visual word recognition.

The BR model predicts similar RTs for words preceded by an unrelated prime regardless of its wordlikeness in lexical decision (see Loth & Davis, 2010b, and Norris & Kinoshita, 2008, for simulations). Consistent with the BR model, all unrelated priming conditions behaved similarly in standard word recognition scenarios (i.e., when the foils were orthographically legal, as in Exper-

Table 4
Mean Correct Response Times (in ms) and Accuracy (in Parentheses) for Words and Nonwords in Experiment 3

Target lexicality	Prime type			
	High-frequency word	Low-frequency word	Pseudoword	Consonant string
Word	644 (96.7)	640 (96.9)	643 (95.6)	647 (95.6)
Pseudoword	701 (97.5)	708 (98.2)	698 (97.9)	709 (97.0)

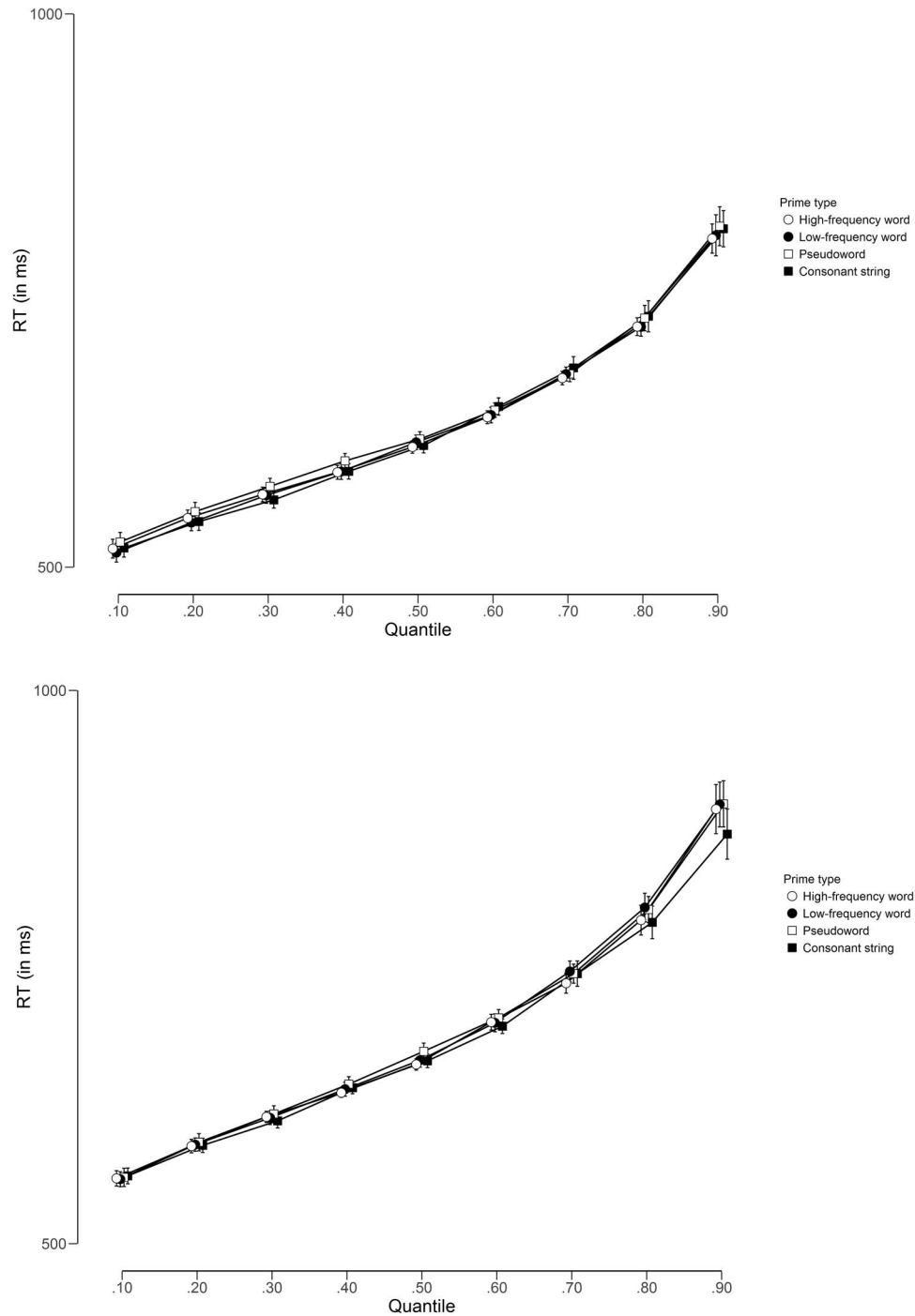


Figure 4. Group reaction time (RT) distributions for the four experimental conditions in Experiment 3 (upper panel: words; lower panel: nonwords) across quantiles (.1 to .9).

iments 1 and 3).³ However, when the word–nonword discrimination was much easier (i.e., using illegal nonwords as foils [e.g., *BEFRM*, *TCFRO*] instead of orthographically legal nonwords), we found an effect from the unrelated primes: Word RTs were longer when the masked prime was a consonant string than when the masked prime was an orthographically legal nonword (*ghdk-TRUE*

³ The only minor differences across conditions were in the accuracy analysis. In Experiment 1, there was a 1.6% increase in the error rates when the primes were composed of random consonants relative to pseudoword primes. In Experiment 3, accuracy was 0.8% higher when the target stimuli were preceded by a word prime than when preceded by a nonword prime.

produced longer response times than *fiok-TRUE*; see Figure 2), and nonword responses were faster when preceded by a consonant-string prime than when preceded by an orthographically legal prime, but only in the initial quantiles of RT distributions (*ghdk-BEFRM* produced faster responses than *fiok-BEFRM*; see Figure 2). The question now is whether this latter finding poses some problems for the BR account of masked priming lexical decision. One might argue that the task in Experiment 2 may have switched from lexical decision to an orthographic legality decision. In an orthographic legality decision task, the BR model could quickly accumulate evidence for “yes” or “no” responses from the masked primes: As all the nonword foils violated the orthographic constraints, then anything legal must be a word. That is, participants could carry over from prime to target the estimate of the probability that the visual input is orthographically well formed, not the lexical status. As a result, the BR model could predict slower “yes” responses when the prime is composed of random consonants than when the prime is an orthographically legal stimulus (e.g., a pseudoword or a word), as actually occurred in Experiment 2. Conversely, the model could predict faster “no” responses when the prime is a consonant string than when the prime is an orthographically legal stimulus (e.g., a pseudoword or a word). Indeed, as indicated in the introduction, the BR model can readily accommodate response congruency effects in two-choice tasks other than lexical decision (e.g., semantic categorization; see de Wit & Kinoshita, 2014). Thus, the present findings offer some empirical support to the BR account of masked priming in standard lexical decision experiments—they are also a demonstration of how the nature of the decision required by the task modulates masked priming effects.

The IA model (McClelland & Rumelhart, 1981) offers different predictions depending on whether there is homogeneous or selective lateral inhibition (see Bayliss et al., 2009, for simulations). Clearly, the homogeneous lateral inhibition hypothesis cannot capture the pattern of effects in any of the three experiments (i.e., this hypothesis wrongly predicts longer “yes” decisions in the high-frequency word priming condition than in the nonword priming conditions). Instead, the pattern of the data with orthographically legal nonwords (Experiments 1 and 3) favors Davis and Lupker’s (2006) selective lateral inhibition hypothesis: The processing of the word *TRUE* is not modified by the wordlikeness of the unrelated prime, whether it is a high-frequency word like *food* or a consonant string like *ghdk*. The only potential drawback is that response times are faster for *food-TRUE* than for *ghdk-TRUE* when the foils were orthographically illegal (Experiment 2). Although the original IA model under the selective lateral inhibition hypothesis cannot predict this difference, one might argue that participants in Experiment 2 were making an orthographic legality judgment rather than a lexical decision. This orthographic legality judgment task would be beyond the scope of the IA model. A somewhat crude approximation would be to use “overall lexical activity” as a source of evidence to speed up responses in this scenario, as in the “fast-guess” criterion proposed by the MROM (Grainger & Jacobs, 1996) or the SCM (see Davis, 2010). If the fast-guess criterion were set low when the nonword foils did not generate practically any lexical activity, those words whose primes produce some lexical activation would enjoy some advantage relative to those words whose primes barely produced any lexical activation. Indeed, Loth and Davis (2010b) conducted simulations that showed that, depending on the value of the parameters, the

SCM could predict an effect of response congruency (e.g., *fiok-TRUE* faster than *ghdk-TRUE*) when the nonword foils are unwordlike and orthographically illegal. However, it is unclear to us how this fast-guess mechanism can accommodate the findings of Experiment 3 with hermit nonwords as foils—note that these stimuli also generate a very low level of activation in the lexicon. Further empirical and computational work is necessary to examine the similarities and differences on word versus nonword responses in lexical decision versus orthographic judgment tasks.

In summary, the present experiments were designed to examine whether response congruency effects with wordlike and unwordlike unrelated primes could be found in lexical decision. When the foils were orthographically legal nonwords, both wordlike and unwordlike priming conditions behaved similarly (Experiments 1 and 3). These findings suggest that participants were unable to establish the lexical status of the masked prime or even to determine the likelihood of its being a word in a standard lexical decision task, thus providing empirical support to the BR model and the selective inhibition hypothesis of the IA model. Furthermore, our findings are a demonstration that detecting invariances (i.e., the null effect of lexical status of unrelated primes on target processing in masked priming lexical decision) is important to progress in science (see Rouder, Speckman, Sun, Morey, & Iverson, 2009). Thus, at a methodological level, the current experiments are useful for choosing the appropriate baseline in masked priming lexical decision experiments. We only found an effect of the unrelated primes on target processing when the foils were orthographically illegal and the primes were composed of consonant strings. However, in this scenario, one might argue that the participants’ responses could have been driven by orthographic properties rather than lexical access. Further experimentation using measures with a better temporal resolution (e.g., ERPs) may offer important insights on the time course of the effects of response congruency across tasks.

References

- Adelman, J. S., Johnson, R. L., McCormick, S. F., McKague, M., Kinoshita, S., Bowers, J. S., . . . Davis, C. J. (2014). A behavioral database for masked form priming. *Behavior Research Methods*, 46, 1052–1067. <http://dx.doi.org/10.3758/s13428-013-0442-y>
- Bates, D., Mächler, M., Bolker, B., & Walker, S. (2015). Fitting linear mixed-effects models using lme4. *Journal of Statistical Software*, 67, 1–48. <http://dx.doi.org/10.18637/jss.v067.i01>
- Bayliss, L., Davis, C. J., Brysbaert, M., Luyten, S., & Rastle, K. (2009). *How are lexical decisions to word targets influenced by unrelated masked primes?* Unpublished manuscript. Retrieved from <http://crr.ugent.be/papers/Bayliss%20et%20al.%202008.pdf>
- Coltheart, M., Davelaar, E., Jonasson, J. T., & Besner, D. (1977). Access to the internal lexicon. In S. Dornic (Ed.), *Attention and performance VI* (pp. 535–555). Hillsdale, NJ: Erlbaum.
- Coltheart, M., Rastle, K., Perry, C., Langdon, R., & Ziegler, J. (2001). DRC: A dual route cascaded model of visual word recognition and reading aloud. *Psychological Review*, 108, 204–256. <http://dx.doi.org/10.1037/0033-295X.108.1.204>
- Davis, C. J. (2010). The spatial coding model of visual word identification. *Psychological Review*, 117, 713–758. <http://dx.doi.org/10.1037/a0019738>
- Davis, C. J., & Lupker, S. J. (2006). Masked inhibitory priming in English: Evidence for lexical inhibition. *Journal of Experimental Psychology: Human Perception and Performance*, 32, 668–687. <http://dx.doi.org/10.1037/0096-1523.32.3.668>

- Dehaene, S., Naccache, L., Cohen, L., Bihan, D. L., Mangin, J. F., Poline, J. B., & Rivière, D. (2001). Cerebral mechanisms of word masking and unconscious repetition priming. *Nature Neuroscience*, 4, 752–758. <http://dx.doi.org/10.1038/89551>
- Dehaene, S., Naccache, L., Le Clec'h, G., Koechlin, E., Mueller, M., Dehaene-Lambertz, G., . . . Le Bihan, D. (1998). Imaging unconscious semantic priming. *Nature*, 395, 597–600. <http://dx.doi.org/10.1038/26967>
- de Wit, B., & Kinoshita, S. (2014). Relatedness proportion effects in semantic categorization: Reconsidering the automatic spreading activation process. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 40, 1733–1744. <http://dx.doi.org/10.1037/xlm0000004>
- Duchon, A., Perea, M., Sebastián-Gallés, N., Martí, A., & Carreiras, M. (2013). EsPal: One-stop shopping for Spanish word properties. *Behavior Research Methods*, 45, 1246–1258. <http://dx.doi.org/10.3758/s13428-013-0326-1>
- Forster, K. I. (1998). The pros and cons of masked priming. *Journal of Psycholinguistic Research*, 27, 203–233. <http://dx.doi.org/10.1023/A:1023202116609>
- Forster, K. I. (2004). Category size effects revisited: Frequency and masked priming effects in semantic categorization. *Brain and Language*, 90, 276–286. [http://dx.doi.org/10.1016/S0093-934X\(03\)00440-1](http://dx.doi.org/10.1016/S0093-934X(03)00440-1)
- Forster, K. I. (2013). How many words can we read at once? More intervenor effects in masked priming. *Journal of Memory and Language*, 69, 563–573. <http://dx.doi.org/10.1016/j.jml.2013.07.004>
- Forster, K. I., & Davis, C. (1984). Repetition priming and frequency attenuation in lexical access. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 10, 680–698. <http://dx.doi.org/10.1037/0278-7393.10.4.680>
- Forster, K. I., & Forster, J. C. (2003). DMDX: A windows display program with millisecond accuracy. *Behavior Research Methods, Instruments & Computers*, 35, 116–124. <http://dx.doi.org/10.3758/BF03195503>
- Forster, K. I., Mohan, K., & Hector, J. (2003). The mechanics of masked priming. In S. Kinoshita & S. Lupker (Eds.), *Masked priming: The state of the art* (pp. 3–37). New York, NY: Psychology Press.
- Gomez, P., Perea, M., & Ratcliff, R. (2013). A diffusion model account of masked versus unmasked priming: Are they qualitatively different? *Journal of Experimental Psychology: Human Perception and Performance*, 39, 1731–1740. <http://dx.doi.org/10.1037/a0032333>
- Grainger, J. (2008). Cracking the orthographic code: An introduction. *Language and Cognitive Processes*, 23, 1–35. <http://dx.doi.org/10.1080/01690960701578013>
- Grainger, J., & Holcomb, P. J. (2009). Watching the word go by: On the time-course of component processes in visual word recognition. *Language and Linguistics Compass*, 3, 128–156. <http://dx.doi.org/10.1111/j.1749-818X.2008.00121.x>
- Grainger, J., & Jacobs, A. M. (1996). Orthographic processing in visual word recognition: A multiple read-out model. *Psychological Review*, 103, 518–565. <http://dx.doi.org/10.1037/0033-295X.103.3.518>
- Huber, D. E., & O'Reilly, R. C. (2003). Persistence and accommodation in short-term priming and other perceptual paradigms: Temporal segregation through synaptic depression. *Cognitive Science*, 27, 403–430. http://dx.doi.org/10.1207/s15516709cog2703_4
- Jacobs, A. M., & Grainger, J. (1992). Testing a semistochastic variant of the interactive activation model in different word recognition experiments. *Journal of Experimental Psychology: Human Perception and Performance*, 18, 1174–1188. <http://dx.doi.org/10.1037/0096-1523.18.4.1174>
- Jacobs, A. M., Grainger, J., & Ferrand, L. (1995). The incremental priming technique: A method for determining within-condition priming effects. *Perception & Psychophysics*, 57, 1101–1110. <http://dx.doi.org/10.3758/BF03208367>
- Keuleers, E., & Brysbaert, M. (2010). Wuggy: A multilingual pseudoword generator. *Behavior Research Methods*, 42, 627–633. <http://dx.doi.org/10.3758/BRM.42.3.627>
- Kinoshita, S., & Hunt, L. (2008). RT distribution analysis of category congruence effects with masked primes. *Memory & Cognition*, 36, 1324–1334. <http://dx.doi.org/10.3758/MC.36.7.1324>
- Kinoshita, S., & Norris, D. (2012). Task-dependent masked priming effects in visual word recognition. *Frontiers in Psychology*, 3, 178. <http://dx.doi.org/10.3389/fpsyg.2012.00178>
- Kuznetsova, A., Brockhoff, P. B., & Christensen, R. H. B. (2016). *lmerTest: Tests for random and fixed effects for linear mixed effect models (lmer objects of lme4 package)*. R package Version 2.0–33. Retrieved from <http://CRAN.Rproject.org/package=lmerTest>
- Lehtonen, M., Monahan, P. J., & Poeppel, D. (2011). Evidence for early morphological decomposition: Combining masked priming with magnetoencephalography. *Journal of Cognitive Neuroscience*, 23, 3366–3379. http://dx.doi.org/10.1162/jocn_a.00035
- Loth, S., & Davis, C. J. (2010a, January). *Response congruency effects in masked primed lexical decision*. Poster presented at the Experimental Psychology Society, London Meeting, London, UK.
- Loth, S., & Davis, C. J. (2010b). Testing computational accounts of response congruency in lexical decision. In E. J. Davelaar (Ed.), *Connectionist models of neurocognition and emergent behavior: From theory to applications* (pp. 173–189). Singapore: World Scientific Publishing. http://dx.doi.org/10.1142/9789814340359_0012
- McClelland, J. L., & Rumelhart, D. E. (1981). An interactive activation model of context effects in letter perception: I. An account of basic findings. *Psychological Review*, 88, 375–407. <http://dx.doi.org/10.1037/0033-295X.88.5.375>
- Naccache, L., & Dehaene, S. (2001). Unconscious semantic priming extends to novel unseen stimuli. *Cognition*, 80, 215–229. [http://dx.doi.org/10.1016/S0010-0277\(00\)00139-6](http://dx.doi.org/10.1016/S0010-0277(00)00139-6)
- Norris, D. (2006). The Bayesian reader: Explaining word recognition as an optimal Bayesian decision process. *Psychological Review*, 113, 327–357. <http://dx.doi.org/10.1037/0033-295X.113.2.327>
- Norris, D. (2009). Putting it all together: A unified account of word recognition and reaction-time distributions. *Psychological Review*, 116, 207–219. <http://dx.doi.org/10.1037/a0014259>
- Norris, D., & Kinoshita, S. (2008). Perception as evidence accumulation and Bayesian inference: Insights from masked priming. *Journal of Experimental Psychology: General*, 137, 434–455. <http://dx.doi.org/10.1037/a0012799>
- Perea, M., Fernández, L., & Rosa, E. (1998). El papel del status léxico y de la frecuencia del estímulo-señal en la condición no relacionada con la técnica de presentación enmascarada del estímulo-señal [The role of the lexical status and word-frequency of the unrelated primes in the masked priming technique]. *Psicológica*, 19, 311–319.
- Perea, M., Gómez, P., & Fraga, I. (2010). Masked nonword repetition effects in yes/no and go/no-go lexical decision: A test of the evidence accumulation and deadline accounts. *Psychonomic Bulletin & Review*, 17, 369–374. <http://dx.doi.org/10.3758/PBR.17.3.369>
- Perea, M., Jiménez, M., & Gómez, P. (2014). A challenging dissociation in masked identity priming with the lexical decision task. *Acta Psychologica*, 148, 130–135. <http://dx.doi.org/10.1016/j.actpsy.2014.01.014>
- Perea, M., Marcet, A., Lozano, M., & Gomez, P. (2018). Is masked priming modulated by memory load? A test of the automaticity of masked identity priming in lexical decision. *Memory & Cognition*. Advance online publication. <http://dx.doi.org/10.3758/s13421-018-0825-5>
- Perea, M., & Rosa, E. (2000). Repetition and form priming interact with neighborhood density at a brief stimulus onset asynchrony. *Psychonomic Bulletin & Review*, 7, 668–677. <http://dx.doi.org/10.3758/BF03213005>
- Ratcliff, R., Gomez, P., & McKoon, G. (2004). A diffusion model account of the lexical decision task. *Psychological Review*, 111, 159–182. <http://dx.doi.org/10.1037/0033-295X.111.1.159>

- R Core Team. (2017). *R: A language and environment for statistical computing*. Vienna, Austria: R Foundation for Statistical Computing. Retrieved from <http://www.R-project.org/>
- Rouder, J. N., Speckman, P. L., Sun, D., Morey, R. D., & Iverson, G. (2009). Bayesian t tests for accepting and rejecting the null hypothesis. *Psychonomic Bulletin & Review*, 16, 225–237. <http://dx.doi.org/10.3758/PBR.16.2.225>
- Sereno, J. A. (1991). Graphemic, associative, and syntactic priming effects at a brief stimulus onset asynchrony in lexical decision and naming. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 17, 459–477. <http://dx.doi.org/10.1037/0278-7393.17.3.459>
- van Heuven, W. J. (2002). Pseudo (2.07) [Computer software]. Retrieved from <http://www.psychology.nottingham.ac.uk/staff/wvh/internal/research/pseudo/index.html>
- van Heuven, W. J., Mandera, P., Keuleers, E., & Brysbaert, M. (2014). SUBTLEX-U. K.: A new and improved word frequency database for British English. *The Quarterly Journal of Experimental Psychology: Human Experimental Psychology*, 67, 1176–1190. <http://dx.doi.org/10.1080/17470218.2013.850521>
- Yarkoni, T., Balota, D., & Yap, M. (2008). Moving beyond Coltheart's N: A new measure of orthographic similarity. *Psychonomic Bulletin & Review*, 15, 971–979. <http://dx.doi.org/10.3758/PBR.15.5.971>

Appendix A

List of Words and Pseudowords in Experiment 1

The stimuli are presented as quintuples: high-frequency word prime, low-frequency word prime, pseudoword prime, consonant string prime, and TARGET.

Word Targets

común, parra, ácofa, bzjms, DISCO; libre, malla, mulmo, lrjrg, NIEVE; serio, labio, gorpo, zglnc, ÁLBUM; gente, diván, ancor, rsmcb, NIETO; ruido, telón, selor, dqplz, ÁNIMO; largo, tenor, nocho, qdfcn, PLAGA; miedo, copla, salla, prdzt, VAGÓN; pobre, huida, lausa, qfnrt, TRIGO; norte, tapiz, mopén, mlztj, RUMBO; dolor, verso, eblas, qftzn, RUBIO; perro, fusil, gipán, zjtfq, FIRMA; carta, rapaz, patuz, crgqz, MONJA; calma, coral, molpa, cjfpm, FRUTA; dulce, beata, becha, mtdbg, LUNAR; deseo, caldo, ralco, cqjfg, CIFRA; tonto, brujo, cueme, sqlnz, FECHA; árbol, polen, ragoz, jfgrd, FERIA; amigo, felpa, sildo, clrbg, MAREA; viejo, tesis, miaje, bsrdm, FALDA; error, nuera, guemo, zfpnj, PLAZO; línea, censo, bunte, pdtrr, COBRE; forma, yegua, llica, rtjzl, GUIÓN; verde, soplo, acube, gsnqm, TEXTO; padre, surco, puaca, rrncg, BOTÓN; trato, sauna, roine, tcbzf, METAL; bella, pinar, émpcr, btdjp, GLOBO; héroe, tapón, mejar, snjgz, RANGO; culpa, toldo, leana, npgzf, MONTE; broma, dorso, bádor, bjdrn, SONTA; calor, tramo, lobio, stqrl, RABIA; vuelo, flora, retio, mbfjz, VACÍO; viaje, olivo, plado, psfdj, ASILO; playa, tinte, farmo, nrlls, METRO; mando, atlas, goror, cdbzt, CICLO; poder, vicio, murco, mrsbc, PACTO; igual, matiz, diepo, psrmf, DEUDA; hotel, galán, ibuaz, lbrnt, PATIO; mitad, busto, anazo, dsglb, CONDE; éxito, forro, nurto, lrdm, ALTAR; lucha, cesta, jeglo, rbcnq, ABEJA; radio, senda, nunto, dptcq, GORRA; listo, prado, menia, gdtlb, COBRA; suelo, ardor, nitar, sflng, PASTA; justo, nogal, zosto, crgbp, POLLO; ayuda, roble, ellir, dpnlz, MARCO; fondo, tacto, cueno, zcrjd, CHINO; chica, guiño, fidro, msqzb, TRONO; mejor, manía, pagle, rnmqb, ACOSO; avión, átomo, gollé, gfpnz, LATÍN; mundo, azote, rusar, dzrjf, ARROZ; bomba, torso, muste, rsbjf, DIOSA; nariz, coste, sabón, fzdmj, LUNES; pared, pajar, zosel, tldgm, GRASA; libro, mango, nagar, qgldr, RUMOR; nivel, brasa, guapo, lgfmd, RENTA; salud, rasgo, larta, cnzmd, PLAZA; lugar, élite, cadmo, jlngp, SUDOR;

juego, bulto, aolta, lgsm, BALÓN; carne, legua, miega, mncgq, CLIMA; mente, chapa, zueno, dfsc, ZORRO; común, parra, ácofa, bzjms, FRENO; libre, malla, mulmo, lrjrg, SUAVE; serio, labio, gorpo, zglnc, OREJA; gente, diván, ancor, rsmcb, DANZA; ruido, telón, selor, dqplz, TRUCO; largo, tenor, nocho, qdfcn, NOBLE; miedo, copla, salla, prdzt, DROGA; pobre, huida, lausa, qfnrt, PLATO; norte, tapiz, mopén, mlztj, BOLSO; dolor, verso, eblas, qftzn, CUEVA; perro, fusil, gipán, zjtfq, PRESO; carta, rapaz, patuz, crgqz, SIGLO; calma, coral, molpa, cjfpm, CURVA; dulce, beata, becha, mtdbg, COPIA; deseo, caldo, ralco, cqjfg, SUCIO; tonto, brujo, cueme, sqlnz, GRADO; árbol, polen, ragoz, jfgrd, BAHÍA; amigo, felpa, sildo, clrbg, INDIO; viejo, tesis, miaje, bsrdm, JAMÓN; error, nuera, guemo, zfpnj, HUECO; línea, censo, bunte, pdtrr, TIGRE; forma, yegua, llica, rtjzl, BRISA; verde, soplo, acube, gsnqm, LÁSER; padre, surco, puaca, rrncg, RIÑÓN; trato, sauna, roine, tcbzf, ESPÍA; bella, pinar, émpcr, btdjp, TURCO; héroe, tapón, mejar, snjgz, GALLO; culpa, toldo, leana, npgzf, MAFIA; broma, dorso, bádor, bjdrn, ORINA; calor, tramo, lobio, stqrl, QUEJA; vuelo, flora, retio, mbfjz, TROZO; viaje, olivo, plado, psfdj, PLENO; playa, tinte, farmo, nrlls, SOLAR; mando, atlas, goror, cdbzt, MÓVIL; poder, vicio, murco, mrsbc, TEMOR; igual, matiz, diepo, psrmf, DUCHA; hotel, galán, ibuaz, lbrnt, MANGA; mitad, busto, anazo, dsglb, BRUJA; éxito, forro, nurto, lrdm, BARRO; lucha, cesta, jeglo, rbcnq, DIETA; radio, senda, nunto, dptcq, CIVIL; listo, prado, menia, gdtlb, FICHA; suelo, ardor, nitar, sflng, MUSEO; justo, nogal, zosto, crgbp, RUEDA; ayuda, roble, ellir, dpnlz, PIANO; fondo, tacto, cueno, zcrjd, TORRE; chica, guiño, fidro, msqzb, CANTO; mejor, manía, pagle, rnmqb, VILLA; avión, átomo, gollé, gfpnz, CAÑÓN; mundo, azote, rusar, dzrjf, GRANO; bomba, torso, muste, rsbjf, LOCAL; nariz, coste, sabón, fzdmj, BORDE; pared, pajar, zosel, tldgm, CUERO; libro, mango, nagar, qgldr, AGUJA; nivel, brasa, guapo, lgfmd, SUSTO; salud, rasgo, larta, cnzmd, FURIA; lugar, élite, cadmo, jlngp, QUESO; juego, bulto, aolta, lgsm, JUDÍO; carne, legua, miega, mncgq, CULTO; mente, chapa, zueno, dfsc, ÓPERA; común, parra, ácofa, bzjms, BINGO; libre, malla, mulmo, lrjrg, TIMÓN; serio, labio, gorpo, zglnc, TARTA;

(Appendices continue)

gente, diván, ancor, rsmcb, FRASE; ruido, telón, selor, dqplz, RUINA; largo, tenor, nocho, qdfcn, FLUJO; miedo, copla, salla, prdzt, BARRA; pobre, huida, lausa, qfnrt, CARRO; norte, tapiz, mopén, mltztj, RATÓN; dolor, verso, eblas, qftzn, ÁCIDO; perro, fusil, gipán, zjtfq, TECHO; carta, rapaz, patuz, crgqz, BURRO; calma, coral, molpa, cjfpm, IDEAL; dulce, beata, becha, mtdbg, OTOÑO; deseo, caldo, ralco, cqjfg, SELVA; tonto, brujo, cueme, sqlnz, PRESA; árbol, polen, ragoz, jfgrd, PRIMA; amigo, felpa, sildo, clrbg, SELLO; viejo, tesis, miaje, bsrdm, AROMA; error, nuera, guemo, zfpnj, CRUCE; línea, censo, bunte, pdtrr, SORDO; forma, yegua, llica, rtjzl, HUESO; verde, soplo, acube, gsnqm, PUNTA; padre, surco, puaca, rmcg, CALVO; trato, sauna, roine, tcbzf, AUTOR; bella, pinar, émper, btdjp, CORTO; héroe, tapón, mejar, snjgz, MOTOR; culpa, toldo, leana, npgzf, FALLO; broma, dorso, bádor, bjdrn, SALTO; calor, tramo, lobio, stqrl, BARCA; vuelo, flora, retio, mbfjz, TENIS; viaje, olivo, plado, psfdj, HACHA; playa, tinte, farmo, nrlls, ATAJO; mando, atlas, goror, cdbzt, UNIÓN; poder, vicio, murco, mrsbc, ENANO; igual, matiz, diepo, psrmf, PLACA; hotel, galán, ibuaz, lbrnt, SOCIO; mitad, busto, anazo, dsgrb, GESTO; éxito, forro, nurto, lrdfm, TABLA; lucha, cesta, jeglo, rbcnq, NIÑEZ; radio, senda, nunto, dptcq, SABOR; listo, prado, menia, gdtlb, ALDEA; suelo, ardor, nitar, sflng, ARAÑA; justo, nogal, zosto, crgbp, APODO; ayuda, roble, ellir, dplnz, VALLE; fondo, tacto, cueno, zcrjd, ACTOR; chica, guiño, fidro, msqzb, VIUDA; mejor, manía, pagle, rnmqb, ÁRABE; avión, átomo, golle, gfpnz, MANTA; mundo, azote, rusar, dzrjf, CREMA; bomba, torso, muste, rsbjf, GUAPA; nariz, coste, sabón, fzdmi, NORMA; pared, pajar, zosel, tldgm, CLAVO; libro, mango, nagar, qgldr, ACERO; nivel, brasa, guepo, lgfmd, DOSIS; salud, rasgo, larta, cnzmd, CIRCO; lugar, élite, cadmo, jlmgp, SUERO; juego, bulto, aolta, lgsrm, POEMA; carne, legua, miega, mncgq, LETRA; mente, chapa, zueno, dfsct, VAPOR; común, parra, ácofa, bzjms, GRITO; libre, malla, mulmo, lrjrg, PLOMO; serio, labio, gorpo, zglnr, RITMO; gente, diván, ancor, rsmcb, PESCA; ruido, telón, selor, dqplz, HORNO; largo, tenor, nocho, qdfcn, BARBA; miedo, copla, salla, prdzt, ETAPA; pobre, huida, lausa, qfnrt, GRIPE; norte, tapiz, mopén, mltztj, SABIO; dolor, verso, eblas, qftzn, JABÓN; perro, fusil, gipán, zjtfq, ACERA; carta, rapaz, patuz, crgqz, SALSA; calma, coral, molpa, cjfpm, LÁPIZ; dulce, beata, becha, mtdbg, CABLE; deseo, caldo, ralco, cqjfg, TRIBU; tonto, brujo, cueme, sqlnz, ÉTICA; árbol, polen, ragoz, jfgrd, CASCO; amigo, felpa, sildo, clrbg, TINTA; viejo, tesis, miaje, bsrdm, CABRA; error, nuera, guemo, zfpnj, TANGO; línea, censo, bunte, pdtnr, JAULA; forma, yegua, llica, rtjzl, BLUSA; verde, soplo, acube, gsnqm, CURSO; padre, surco, puaca, rmcg, LABOR; trato, sauna, roine, tcbzf, SIGNO; bella, pinar, émper, btdjp, RECTA; héroe, tapón, mejar, snjgz, CELDA; culpa, toldo, leana, npgzf, BOTÍN; broma, dorso, bádor, bjdrn, VÍDEO; calor, tramo, lobio, stqrl, LICOR; vuelo, flora, retio, mbfjz, GAFAS; viaje, olivo, plado, psfdj, DUELO; playa, tinte, farmo, nrlls, TÚNEL; mando, atlas, goror, cdbzt, DUEÑA; poder, vicio, murco, mrsbc, OVEJA; igual, matiz, diepo, psrmf, BUZÓN; hotel, galán, ibuaz, lbrnt, CRUEL;

mitad, busto, anazo, dsgrb, MORAL; éxito, forro, nurto, lrdfm, PLANO; lucha, cesta, jeglo, rbcnq, ROBOT; radio, senda, nunto, dptcq, VARÓN; listo, prado, menia, gdtlb, VENTA; suelo, ardor, nitar, sflng, BICHO; justo, nogal, zosto, crgbp, TROPA; ayuda, roble, ellir, dplnz, AVISO; fondo, tacto, cueno, zcrjd, MOSCA; chica, guiño, fidro, msqzb, CAJÓN; mejor, manía, pagle, rnmqb, TAREA; avión, átomo, golle, gfpnz, TORTA; mundo, azote, rusar, dzrjf, PLUMA; bomba, torso, muste, rsbjf, DRAMA; nariz, coste, sabón, fzdmi, RIVAL; pared, pajar, zosel, tldgm, VALLA; libro, mango, nagar, qgldr, PARTO; nivel, brasa, guepo, lgfmd, PULSO; salud, rasgo, larta, cnzmd, DUQUE; lugar, élite, cadmo, jlmgp, DÉBIL; juego, bulto, aolta, lgsrm, BANDO; carne, legua, miega, mncgq, HUEVO; mente, chapa, zueno, dfsct, TALLA.

Pseudoword Targets

clase, suela, firia, bzrnd, CUODA; coche, boina, vomón, lbctr, BREMO; santo, casta, mucor, nbmlf, AGACO; joven, gasto, ruena, nzftr, BLIGA; prisa, saldo, drolo, grjds, TUCOR; medio, secta, fluja, lsgfd, SASGA; buena, jarra, pasno, jdgrf, SECIS; fuego, farol, goren, gfrdm, MAPTO; grupo, verbo, badón, fnbgs, FLERO; llave, andén, bejua, mgqbn, RUPTA; sitio, cauce, arnia, rlfnd, DIRSA; falta, faena, fazón, dpqtm, URIÁN; noche, grifo, legor, mrbtf, SIFLE; feliz, gruta, tabez, tzpgl, CASPO; mujer, limbo, gosál, tfcrr, OCEGO; calle, aleta, maver, rbltj, ÉCITO; golpe, peine, pudes, sfbm, TABIA; único, fobia, filza, scrzg, AVURO; pelea, útero, bomal, rgncd, RALTA; brazo, manto, cergo, ldgrc, BANJA; razón, cojín, broga, mrbnr, SOEGO; cielo, cisne, mepla, mnrbj, PLARO; negro, mareo, ralmo, zlngt, PRACO; reloj, cloro, runco, crzrf, ACUFO; barco, pavor, orazo, rbsdl, SUSNE; color, ancla, valbo, fgrrd, RULLO; final, muslo, prana, tnbms, HUIBO; vista, tacón, lorio, srfrn, DITAL; mayor, mueca, baero, qpnjf, DRAÑO; sueño, ámbar, croja, cmrdt, SOLLA; fácil, cuota, delle, zqbmc, BRANA; orden, sodio, muvor, dcpng, NURJA; causa, fogón, plame, drcpf, LLUPO; hielo, trama, mocer, dmqbp, TARBO; oeste, fiera, prujo, dslmn, COSCE; trago, rigor, ficle, lfrqs, SUCEA; traje, cutis, pusto, trfbd, CHIGU; baile, vejez, mante, stldb, SUNOR; señal, pudor, faubo, lfrnt, MOPRA; total, pompa, rargo, gzdrn, MENZA; señor, musgo, vimal, dngrc, CRENO; novia, tripa, carre, qrcmg, SENCO; punto, fauna, mobun, mgcjt, NIJÓN; nuevo, pasto, relío, rmlmf, MATRA; honor, fibra, zuche, qdjzp, PURCA; banda, prosa, fánel, tqjdc, FÁMER; doble, folio, erdin, sqcfm, ÁLFEM; gusto, miope, mogre, fpzrj, RAPIO; lleno, momia, gopia, tpbqg, GRUBA; parte, vigor, miaco, jsllf, VIASA; madre, tórax, cunra, rbnjg, CEGRO; hogar, lirio, mapre, nzrsf, RUJÓN; valor, rosál, lorna, mrbqt, DIEME; capaz, bambú, vaípe, drnzj, BIFÓN; banco, ansia, úmaso, brqmt, NACEO; corte, dardo, trida, nlrqt, ROTÚN; campo, trapo, ruezo, lgcrd, JARSA; reina, palmo, meper, sldfp, ASENSA; leche, funda, nonte, mdtbj, AIDOR; papel, fresa, leluz, lrsqg, NUEJO; clase, suela, firia, bzrnd, REULA; coche, boina, vomón, lbctr, HERRU; santo, casta, mucor, nbmlf, PURTA; joven, gasto, ruena, nzftr, IBEOL; prisa, saldo, drolo, grjds, RULGO; medio, secta, fluja, lsgfd, LUTIA;

(Appendices continue)

buena, jarra, pasno, jdghz, RAPOR; fuego, farol, goren, gfrdm, HAUSO; grupo, verbo, badón, fnbgs, SIZAL; llave, andén, bejua, mgqbn, DONTA; sitio, cauce, arnia, rlfnd, ROBUC; falta, faena, fazón, dpqtm, ÁVADA; noche, grifo, legor, mrbt, LLITO; feliz, gruta, tabez, tzpgl, CRUSA; mujer, limbo, gosál, tferr, VECHA; calle, aleta, maver, rbltj, FLENA; golpe, peine, pudes, sbrn, MABÓN; único, fobia, filza, scrzg, DARJA; pelea, útero, bomal, rgncd, LANCA; brazo, manto, cergo, ldgrc, COVIO; razón, cojín, broga, mrbnr, DURCO; cielo, cisne, mepla, mnrjb, GAPÓN; negro, mareo, ralmo, zlngt, FLADO; reloj, cloro, runco, crzrf, CURNÓ; barco, pavor, orazo, rbsdl, GODIS; color, ancla, valbo, fgrrd, VAHÚA; final, muslo, prana, tnbs, QUIÓS; vista, tacón, lorio, srfnm, BONTA; mayor, mueca, baero, qpnjf, SUCNO; sueño, ámbra, croja, cmrdt, PUSGO; fácil, cuota, delle, zqbmc, RIERO; orden, sodio, muvor, dcpng, GOGRE; causa, fogón, plame, drcpf, BEDRA; hielo, trama, mocer, dmqbp, BURSO; oeste, fiero, prujo, dslmn, DACHO; trago, rigor, ficle, lfrqs, TIFRO; traje, cutis, puesto, trfbd, PEISO; baile, vejez, mante, stldb, FIDOR; señal, pudor, faubo, lfrnt, GUEVA; total, pompa, rargo, gzdrn, ECINÁ; señor, musgo, vimal, dngtc, DOCHO; novia, tripa, carre, qrcmg, ACOFO; punto, fauna, mobun, mgcjt, SÚNIL; nuevo, pasto, relío, rnlmf, GIJEZ; honor, fibra, zuche, qdjz, TELZA; banda, prosa, fánel, tqjdc, DUVIO; doble, folio, erdin, sqcfm, CHOGO; gusto, miope, mogre, fpzrj, GERTA; lleno, momia, gopia, tpbqg, BARGO; parte, vigor, miaco, jsrlf, GAULO; madre, tórax, cunra, rbnjg, GULLA; hogar, lirio, mapre, nzrsf, TRUMO; valor, rosál, lorna, mrbqt, BRADA; capaz, bambú, vaípe, drnzj, PUCES; banco, ansia, úmaso, brqmt, RERTA; corte, dardo, trida, nlrqt, OCATA; campo, trapo, ruego, lgcrd, PRAMA; reina, palmo, meper, sldfp, RAUGA; leche, funda, nonte, mdtbj, MAFLE; papel, fresa, leluz, lrsqg, DARRO; clase, suela, firia, bzrnd, GAÑOS; coche, boina, vomón, lbctr, BRIME; santo, casta, mucor, nbmlf, ONRÍO; joven, gasto, ruena, nzftr, CODOR; prisa, saldo, drolo, grjds, ANCEO; medio, secta, fluja, lsgfd, ÁMAJO; buena, jarra, pasno, jdghz, PLAJO; fuego, farol, goren, gfrdm, CHUDO; grupo, verbo, badón, fnbgs, BACHA; llave, andén, bejua, mgqbn, FOSTO; sitio, cauce, arnia, rlfnd, FÁJIZ; falta, faena, fazón, dpqtm, ACHAZ; noche, grifo, legor, mrbt, CAUPA; feliz, gruta, tabez, tzpgl, PLUNE; mujer, limbo, gosál, tferr, CALSO; calle, aleta, maver, rbltj, ADETO; golpe, peine, pudes, sbrn, MOTRE; único, fobia, filza, scrzg, CISON; pelea, útero, bomal, rgncd, LANGO; brazo, manto, cergo, ldgrc, RINTA; razón, cojín, broga, mrbnr, ECIRO; cielo, cisne, mepla, mnrjb, TAPTO; negro, mareo, ralmo, zlngt, LEBRA; reloj, cloro, runco, crzrf, ERTAR; barco, pavor, orazo, rbsdl, GADÓN; color, ancla, valbo, fgrrd, RABRA; final, muslo, prana, tnbs, LARÓN; vista, tacón, lorio, srfnm, CLASA; mayor, mueca, baero, qpnjf, ACIZA; sueño, ámbra, croja, cmrdt, CIPEL; fácil, cuota, delle, zqbmc, MESTE; orden, sodio, muvor, dcpng, CHEGO; causa, fogón, plame, drcpf,

LICHA; hielo, trama, mocer, dmqbp, ACIGO; oeste, fiero, prujo, dslmn, CHOJA; trago, rigor, ficle, lfrqs, CRIAL; traje, cutis, puesto, trfbd, FLUCA; baile, vejez, mante, stldb, BACHO; señal, pudor, faubo, lfrnt, FLOGO; total, pompa, rargo, gzdrn, OCOZA; señor, musgo, vimal, dngtc, GAUTO; novia, tripa, carre, qrcmg, BRIJE; punto, fauna, mobun, mgcjt, QUIVO; nuevo, pasto, relío, rnlmf, PREMA; honor, fibra, zuche, qdjz, JANCO; banda, prosa, fánel, tqjdc, SACEO; doble, folio, erdin, sqcfm, TEDER; gusto, miope, mogre, fpzrj, ENDIA; lleno, momia, gopia, tpbqg, CUNZA; parte, vigor, miaco, jsrlf, FARRE; madre, tórax, cunra, rbnjg, FADÍO; hogar, lirio, mapre, nzrsf, DRAJA; valor, rosál, lorna, mrbqt, GONTO; capaz, bambú, vaípe, drnzj, BONCE; banco, ansia, úmaso, brqmt, OCENO; corte, dardo, trida, nlrqt, DEMIO; campo, trapo, ruego, lgcrd, ACINE; reina, palmo, meper, sldfp, DUNTO; leche, funda, nonte, mdtbj, NULCA; papel, fresa, leluz, lrsqg, SATRA; clase, suela, firia, bzrnd, LIRRO; coche, boina, vomón, lbctr, ACEFO; santo, casta, mucor, nbmlf, RINCO; joven, gasto, ruena, nzftr, PAMIA; prisa, saldo, drolo, grjds, GODÓN; medio, secta, fluja, lsgfd, TORAR; buena, jarra, pasno, jdghz, NICEA; fuego, farol, goren, gfrdm, JUIRA; grupo, verbo, badón, fnbgs, BAGOR; llave, andén, bejua, mgqbn, GUCHE; sitio, cauce, arnia, rlfnd, ADOCA; falta, faena, fazón, dpqtm, GRESO; noche, grifo, legor, mrbt, GORNA; feliz, gruta, tabez, tzpgl, RARTO; mujer, limbo, gosál, tferr, DABOR; calle, aleta, maver, rbltj, PEGRO; golpe, peine, pudes, sbrn, SUAJA; único, fobia, filza, scrzg, LLASO; pelea, útero, bomal, rgncd, VALTO; brazo, manto, cergo, ldgrc, ERTOR; razón, cojín, broga, mrbnr, CÉTOL; cielo, cisne, mepla, mnrjb, OCEZO; negro, mareo, ralmo, zlngt, JOLLO; reloj, cloro, runco, crzrf, POMAL; barco, pavor, orazo, rbsdl, SASIO; color, ancla, valbo, fgrrd, GACÓN; final, muslo, prana, tnbs, RUZMO; vista, tacón, lorio, srfnm, FESDA; mayor, mueca, baero, qpnjf, CEGLA; sueño, ámbra, croja, cmrdt, CEPTO; fácil, cuota, delle, zqbmc, OMIZA; orden, sodio, muvor, dcpng, PUCAL; causa, fogón, plame, drcpf, NURDO; hielo, trama, mocer, dmqbp, HOMO; oeste, fiero, prujo, dslmn, DRUMA; trago, rigor, ficle, lfrqs, VELDA; traje, cutis, puesto, trfbd, BLETO; baile, vejez, mante, stldb, CHODO; señal, pudor, faubo, lfrnt, LENIA; total, pompa, rargo, gzdrn, VUNDA; señor, musgo, vimal, dngtc, GAUTA; novia, tripa, carre, qrcmg, BROPO; punto, fauna, mobun, mgcjt, MUSTO; nuevo, pasto, relío, rnlmf, RORGO; honor, fibra, zuche, qdjz, MÓPEL; banda, prosa, fánel, tqjdc, LARÍN; doble, folio, erdin, sqcfm, NULLO; gusto, miope, mogre, fpzrj, GUBIO; lleno, momia, gopia, tpbqg, DABÓN; parte, vigor, miaco, jsrlf, LAFIO; madre, tórax, cunra, rbnjg, RECHO; hogar, lirio, mapre, nzrsf, BOMPA; valor, rosál, lorna, mrbqt, ZUBÍO; capaz, bambú, vaípe, drnzj, VABRO; banco, ansia, úmaso, brqmt, CAINO; corte, dardo, trida, nlrqt, NANÓN; campo, trapo, ruego, lgcrd, FRIMA; reina, palmo, meper, sldfp, DARRA; leche, funda, nonte, mdtbj, BURAR; papel, fresa, leluz, lrsqg, CLIGA

(Appendices continue)

Appendix B

List of Nonword Targets In Experiment 2

The stimuli are presented as quintuples: high-frequency word prime, low-frequency word prime, pseudoword prime, consonant string prime, and TARGET. (The word trials are the same as in Experiment 1).

Nonword Targets

clase, suela, firia, bznrd, CBFDA; coche, boina, vomón, lbctr, BEDMT; santo, casta, mucor, nbmlf, AGBCO; joven, gasto, ruena, nzftr, BLIGC; prisa, saldo, drolo, grjds, TKCOR; medio, secta, fluja, lsgfd, SASGP; buena, jarra, pasno, jdgfz, SDCIS; fuego, farol, goren, gfrdm, MAPLR; grupo, verbo, badón, fnbgs, FLBTO; llave, andén, bejua, mgqbn, RUPTP; sitio, cauce, arnia, rlfnd, DMRSA; falta, faena, fazón, dpqtm, URDCN; noche, grifo, legor, mrbtf, SGFTE; feliz, gruta, tabez, tzpgl, CATPB; mujer, limbo, gosál, tfcrr, PCTGO; calle, aleta, maver, rbltj, ÉCSCP; golpe, peine, pudes, sfbrn, TPBIA; único, fobia, filza, scrzg, ALNDT; pelea, útero, bomal, rgncd, RJLTA; brazo, manto, cergo, ldgrc, BANJZ; razón, cojín, broga, mrbnr, SGPGO; cielo, cisne, mepla, mnrbj, PLARDP; negro, mareo, ralmo, zlngt, PRBCO; reloj, cloro, runco, crzrf, ACRFD; barco, pavor, orazo, rbsdl, SDTNE; color, ancla, valbo, fgrrd, RUBPC; final, muslo, prana, tnbms, HJPBO; vista, tacón, lorio, srfnm, DITRL; mayor, mueca, baero, qpnjf, DRTÑO; sueño, ámbar, croja, cmrdt, SOLBF; fácil, cuota, delle, zqbmc, TNRGA; orden, sodio, muvor, dcpng, NURKJ; causa, fogón, plame, drcpf, LLPMO; hielo, trama, mocer, dmqbp, TABRZ; oeste, fierá, prujo, dslmn, CFSCE; trago, rigor, ficle, lfrqs, SUCMR; traje, cutis, pusto, trfbd, CHJGU; baile, vejez, mante, stldb, DUBMR; señal, pudor, faubo, lfrnt, MGPRÁ; total, pompa, rargo, gzdrn, METSL; señor, musgo, vimal, dngtc, CRPNO; novia, tripa, carre, qrcmg, SENTB; punto, fauna, mobun, mgcjt, NMJÓN; nuevo, pasto, relío, rnlmf, MASRP; honor, fibra, zuche, qdjzp, PBTCA; banda, prosa, fánel, tqjdc, FAPMT; doble, folio, erdin, sqcfm, DBTEM; gusto, miope, mogre, fpzrj, RAPGC; lleno, momia, gopia, tpbqg, GRTBA; parte, vigor, miaco, jsrlf, VIMTF; madre, tórax, cunra, rbnjg, CJGRO; hogar, lirio, mapre, nzrsf, RULJM; valor, rosál, lorna, mrbqt, DTPME; capaz, bambú, vaipe, drnzj, BIFCN; banco, ansia, úmaso, brqmt, NMCUO; corte, dardo, trida, nlrqt, ROTNL; campo, trapo, ruezo, lgrcd, JTRSA; reina, palmo, meper, sldfp, ASPGT; leche, funda, nonte, mdthj, DLTOR; papel, fresa, leluz, lrsqg, MUDLD; clase, suela, firia, bznrd, RTHLA; coche, boina, vomón, lbctr, HEFNR; santo, casta, mucor, nbmlf, TMSRA; joven, gasto, ruena, nzftr, IBSTP; prisa, saldo, drolo, grjds, RBLGO; medio, secta, fluja, lsgfd, LUTGF; buena, jarra, pasno, jdgfz, RTBOR; fuego, farol, goren, gfrdm, HASFZ; grupo, verbo, badón, fnbgs, SBZAL; llave, andén, bejua, mgqbn, DOBTC; sitio, cauce, arnia, rlfnd, RTBUB; falta, faena, fazón, dpqtm, ÁVNTP; noche, grifo, legor, mrbtf, LCBTO; feliz, gruta, tabez, tzpgl, CUBBJ; mujer, limbo, gosál, tfcrr, TMHCA; calle, aleta, maver, rbltj, FLGEM; golpe, peine, pudes, sfbrn, MGVÓN; único, fobia, filza, scrzg, DARJC; pelea, útero, bomal, rgncd,

LBNTA; brazo, manto, cergo, ldgrc, COVBP; razón, cojín, broga, mrbnr, DTRCO; cielo, cisne, mepla, mnrbj, GAPBN; negro, mareo, ralmo, zlngt, FBPDO; reloj, cloro, runco, crzrf, CUBNM; barco, pavor, orazo, rbsdl, GLDIT; color, ancla, valbo, fgrrd, VAHPT; final, muslo, prana, tnbms, QBJÓS; vista, tacón, lorio, srfnm, BOGTC; mayor, mueca, baero, qpnjf, SDCNO; sueño, ámbar, croja, cmrdt, PUSGF; fácil, cuota, delle, zqbmc, RITRB; orden, sodio, muvor, dcpng, GCTRE; causa, fogón, plame, drcpf, BEFRM; hielo, trama, mocer, dmqbp, PMRSO; oeste, fierá, prujo, dslmn, DACHG; trago, rigor, ficle, lfrqs, TCFRO; traje, cutis, pusto, trfbd, PEDSC; baile, vejez, mante, stldb, FBDOR; señal, pudor, faubo, lfrnt, GUBCF; total, pompa, rargo, gzdrn, TCSÑA; señor, musgo, vimal, dngtc, DOCGP; novia, tripa, carre, qrcmg, GCPFO; punto, fauna, mobun, mgcjt, SÚCBL; nuevo, pasto, relío, rnlmf, GCJEZ; honor, fibra, zuche, qdjzp, TELZP; banda, prosa, fánel, tqjdc, GLVIO; doble, folio, erdin, sqcfm, CORGC; gusto, miope, mogre, fpzrj, GCRTA; lleno, momia, gopia, tpbqg, BARGP; parte, vigor, miaco, jsrlf, GFTLO; madre, tórax, cunra, rbnjg, GULBZ; hogar, lirio, mapre, nzrsf, TRBMO; valor, rosál, lorna, mrbqt, BAPDP; capaz, bambú, vaipe, drnzj, PCMES; banco, ansia, úmaso, brqmt, RERTG; corte, dardo, trida, nlrqt, PCBTA; campo, trapo, ruezo, lgrcd, PGAMF; reina, palmo, meper, sldfp, RADGB; leche, funda, nonte, mdthj, MCFTE; papel, fresa, leluz, lrsqg, DATRD; clase, suela, firia, bznrd, GFCUN; coche, boina, vomón, lbctr, BIRTR; santo, casta, mucor, nbmlf, LPTÍO; joven, gasto, ruena, nzftr, CODMR; prisa, saldo, drolo, grjds, ZNPEO; medio, secta, fluja, lsgfd, ÁMKJP; buena, jarra, pasno, jdgfz, CLTJO; fuego, farol, goren, gfrdm, CUFSL; grupo, verbo, badón, fnbgs, TKCHA; llave, andén, bejua, mgqbn, FUSTM; sitio, cauce, arnia, rlfnd, FKJIZ; falta, faena, fazón, dpqtm, ACHBT; noche, grifo, legor, mrbtf, CLTPA; feliz, gruta, tabez, tzpgl, PUCLM; mujer, limbo, gosál, tfcrr, SCLSO; calle, aleta, maver, rbltj, AEDTB; golpe, peine, pudes, sfbrn, MPTBE; único, fobia, filza, scrzg, CITFN; pelea, útero, bomal, rgncd, NLPPO; brazo, manto, cergo, ldgrc, RITNF; razón, cojín, broga, mrbnr, PBTRO; cielo, cisne, mepla, mnrbj, TAPTL; negro, mareo, ralmo, zlngt, LSRGA; reloj, cloro, runco, crzrf, ERTGR; barco, pavor, orazo, rbsdl, GCDÓN; color, ancla, valbo, fgrrd, MUBTL; final, muslo, prana, tnbms, ZLRÓN; vista, tacón, lorio, srfnm, CETLR; mayor, mueca, baero, qpnjf, GCTZA; sueño, ámbar, croja, cmrdt, CIPBL; fácil, cuota, delle, zqbmc, MFSTE; orden, sodio, muvor, dcpng, CEHPG; causa, fogón, plame, drcpf, LZHCA; hielo, trama, mocer, dmqbp, AIBGZ; oeste, fierá, prujo, dslmn, CHPJA; trago, rigor, ficle, lfrqs, CIRLB; traje, cutis, pusto, trfbd, FLTCA; baile, vejez, mante, stldb, BACDP; señal, pudor, faubo, lfrnt, FLCGO; total, pompa, rargo, gzdrn, OCPZT; señor, musgo, vimal, dngtc, TLGTO; novia, tripa, carre, qrcmg, BILJP; punto, fauna, mobun, mgcjt, TPMVO; nuevo, pasto, relío, rnlmf, PECCG; honor, fibra, zuche, qdjzp, JTNCO; banda, prosa, fánel, tqjdc, SAFDV; doble, folio,

(Appendices continue)

erdin, sqcfm, LMTER; gusto, miope, mogre, fpzrj, EDNPM; lleno, momia, gopia, tpbqg, RNZPA; parte, vigor, miaco, jsrlf, FABST; madre, tórax, cunra, rbnjg, FGDÍO; hogar, lirio, mapre, nzrsf, DABMG; valor, rosál, lorna, mrbqt, PGNTQ; capaz, bambú, vaipe, drnzj, BODCT; banco, ansia, úmaso, brqmt, RTCNO; corte, dardo, trida, nlrqt, DEMPL; campo, trapo, ruezo, lgcrd, PCLÑE; reina, palmo, meper, sldfp, DUNTÑ; leche, funda, nonte, mdthj, NLGCA; papel, fresa, leluz, lrsqg, SAPGC; clase, suela, firia, bznrd, LPSRO; coche, boina, vomón, lbctr, AEDFP; santo, casta, mucor, nbmlf, RMNCO; joven, gasto, ruena, nzftr, PAMGB; prisa, saldo, drolo, grjds, GBDÓN; medio, secta, fluja, lsgfd, TODPC; buena, jarra, pasno, jdgfz, CPCEA; fuego, farol, goren, gfrdm, JUCBT; grupo, verbo, badón, fnbgs, BTVOR; llave, andén, bejua, mgqbn, GUCPL; sitio, cauce, arnia, rlfnd, SRDCA; falta, faena, fazón, dpqtm, GEDSP; noche, grifo, legor, mrbtf, GTRNA; feliz, gruta, tabez, tzpgl, RARST; mujer, limbo, gosál, tfcrr, DSVOR; calle, aleta, maver, rbltj, PEGRD; golpe, peine, pudes, sbrn, SBTJÁ; único, fobia, filza, scrzg, LAVZS; pelea, útero, bomal, rgncd, BTLTO; brazo, manto, cergo, ldgrc, ERTBR; razón, cojín, broga, mrbnr, CPTOL; cielo, cisne, mepla, mnrbj, OECZP; negro, mareo, ralmo, zlngt, JCLPO; reloj, cloro, runco, crrzf, POMCL;

barco, pavor, orazo, rbsdl, SGSIO; color, ancla, valbo, fgrrd, GÁCTN; final, muslo, prana, tnbms, RBZMO; vista, tacón, lorio, srfrn, FESDC; mayor, mueca, baero, qpnjf, CJGLA; sueño, ámba, croja, cmrdt, CEPTZ; fácil, cuota, delle, zqbmc, JMRZA; orden, sodio, muvor, dcpng, PUVRL; causa, fogón, plame, drcpf, NJRDO; hielo, trama, mocer, dmqbp, HOBTC; oeste, fierá, prujo, dslmn, DPTMA; trago, rigor, ficle, lfrqs, VEPTD; traje, cutis, pusto, trfbd, BLCTO; baile, vejez, mante, stldb, COFDD; señal, pudor, faubo, lfrnt, LPNIA; total, pompa, rargo, gzdrn, VUNLR; señor, musgo, vimal, dngtc, GDHTA; novia, tripa, carre, qrcmg, BOFBG; punto, fauna, mobun, mgcjt, MGRTO; nuevo, pasto, relío, rnlmf, ROFBT; honor, fibra, zuche, qdjz, MTEPL; banda, prosa, fánel, tqjdc, LACLN; doble, folio, erdin, sqcfm, NPLBO; gusto, miope, mogre, fpzrj, BUBPC; lleno, momia, gopia, tpbqg, TFBÓZ; parte, vigor, miaco, jsrlf, SAFLR; madre, tórax, cunra, rbnjg, RMCHO; hogar, lirio, mapre, nzrsf, BOGPC; valor, rosál, lorna, mrbqt, ZRBÍO; capaz, bambú, vaipe, drnzj, ZANBL; banco, ansia, úmaso, brqmt, CPCNO; corte, dardo, trida, nlrqt, NANLR; campo, trapo, ruezo, lgcrd, FPÑMA; reina, palmo, meper, sldfp, DAGLP; leche, funda, nonte, mdthj, BDTER; papel, fresa, leluz, lrsqg, CILGD

Appendix C

List of Nonword Targets in Experiment 3

The stimuli are presented as quintuples: high-frequency word prime, low-frequency word prime, pseudoword prime, consonant string prime, and TARGET. (The word trials are the same as in Experiments 1 and 2).

Nonword Targets

clase, suela, firia, bznrd, ÁCILI; coche, boina, vomón, lbctr, ADUBA; santo, casta, mucor, nbmlf, AECAZ; joven, gasto, ruena, nzftr, AGESO; prisa, saldo, drolo, grjds, ÁGEZO; medio, secta, fluja, lsgfd, AHEVA; buena, jarra, pasno, jdgfz, ALOVO; fuego, farol, goren, gfrdm, APANÚ; grupo, verbo, badón, fnbgs, ASNIP; llave, andén, bejua, mgqbn, ÁVECO; sitio, cauce, arnia, rlfnd, AYEHO; falta, faena, fazón, dpqtm, AZANI; noche, grifo, legor, mrbtf, AZOXA; feliz, gruta, tabez, tzpgl, AZUGO; mujer, limbo, gosál, tfcrr, BAPEI; calle, aleta, maver, rbltj, BESCE; golpe, peine, pudes, sbrn, BIAPA; único, fobia, filza, scrzg, BIFRE; pelea, útero, bomal, rgncd, BIRFO; brazo, manto, cergo, ldgrc, BILTA; razón, cojín, broga, mrbnr, BIODI; cielo, cisne, mepla, mnrbj, BISNA; negro, mareo, ralmo, zlngt, BITIL; reloj, cloro, runco, crrzf, BOGÚN; barco, pavor, orazo, rbsdl, BÓPAZ; color, ancla, valbo, fgrrd, BREMI; final, muslo, prana, tnbms, BRUPE; vista, tacón, lorio, srfrn, BUMIL; mayor, mueca, baero, qpnjf, BUNTE; sueño, ámba, croja, cmrdt, BUPOE; fácil, cuota, delle, zqbmc, CAXIL; orden, sodio, muvor, dcpng, CECÓN; causa, fogón, plame, drcpf, CEFRI; hielo, trama, mocer, dmqbp, CELUR; oeste, fierá, prujo, dslmn, CILBE; trago, rigor, ficle, lfrqs, CLELE;

traje, cutis, pusto, trfbd, CLUFE; baile, vejez, mante, stldb, COLCI; señal, pudor, faubo, lfrnt, CRAMS; total, pompa, rargo, gzdrn, CUFER; señor, musgo, vimal, dngtc, CUPAC; novia, tripa, carre, qrcmg, DAENO; punto, fauna, mobun, mgcjt, DIBEO; nuevo, pasto, relío, rnlmf, DILCI; honor, fibra, zuche, qdjz, DOBUR; banda, prosa, fánel, tqjdc, DOCHI; doble, folio, erdin, sqcfm, DRUFI; gusto, miope, mogre, fpzrj, DUIPÁ; lleno, momia, gopia, tpbqg, DUSPE; parte, vigor, miaco, jsrlf, ECAZO; madre, tórax, cunra, rbnjg, ÉGACA; hogar, lirio, mapre, nzrsf, ELGÓN; valor, rosál, lorna, mrbqt, EÑEBA; capaz, bambú, vaipe, drnzj, EÑOZO; banco, ansia, úmaso, brqmt, EPEJA; corte, dardo, trida, nlrqt, ERTOL; campo, trapo, ruezo, lgcrd, ÉSOCO; reina, palmo, meper, sldfp, ÉTOXA; leche, funda, nonte, mdthj, ETUMA; papel, fresa, leluz, lrsqg, EUCOR; clase, suela, firia, bznrd, EVEÓN; coche, boina, vomón, lbctr, EVUMA; santo, casta, mucor, nbmlf, FAIBA; joven, gasto, ruena, nzftr, FALGI; prisa, saldo, drolo, grjds, FAGPE; medio, secta, fluja, lsgfd, FANIR; buena, jarra, pasno, jdgfz, FÉFAR; fuego, farol, goren, gfrdm, FEGUZ; grupo, verbo, badón, fnbgs, FESMA; llave, andén, bejua, mgqbn, FLAFE; sitio, cauce, arnia, rlfnd, FLEFI; falta, faena, fazón, dpqtm, FLONI; noche, grifo, legor, mrbtf, FOLFE; feliz, gruta, tabez, tzpgl, FOLMU; mujer, limbo, gosál, tfcrr, FOLUZ; calle, aleta, maver, rbltj, FOMOR; golpe, peine, pudes, sbrn, FOSDI; único, fobia, filza, scrzg, FROXA; pelea, útero, bomal, rgncd, FUEVI; brazo, manto, cergo, ldgrc, FULPI; razón, cojín, broga, mrbnr, GALSI; cielo, cisne, mepla, mnrbj, GAZUR; negro, mareo, ralmo,

(Appendices continue)

zlngt, GEAJA; reloj, cloro, runco, crrzf, GEFÚN; barco, pavor, orazo, rbsdl, GERFO; color, ancla, valbo, fgrrd, GILZO; final, muslo, prana, tnbms, GLEME; vista, tacón, lorio, srfrnm, GOPIS; mayor, mueca, baero, qpnjf, GUGOZ; sueño, ámbra, croja, cmrdt, GUNGO; fácil, cuota, delle, zqbmc, GUNTI; orden, sodio, muvor, dcpng, GUPÉN; causa, fogón, plame, drcpf, GUSNA; hielo, trama, mocer, dmqbp, HANAC; oeste, fieras, prujo, dslmn, HARUR; trago, rigor, ficle, lfrqs, HECRE; traje, cutis, pusto, trfbd, HESNE; baile, vejez, mante, stldb, HESUR; señal, pudor, faubo, lfrnt, HIBUR; total, pompa, rargo, gzdrrn, HOLTO; señor, musgo, vimal, dngtc, HUOTE; novia, tripa, carre, qrcmg, HUSCE; punto, fauna, mobun, mgcjt, IBAFE; nuevo, pasto, relío, rnlmf, ILAPÚ; honor, fibra, zuche, qdjzp, ILASI; banda, prosa, fánel, tqjdc, INENO; doble, folio, erdin, sqcfm, ISGER; gusto, miope, mogre, fpzrj, ITOBO; lleno, momia, gopia, tpbqg, IXTAR; parte, vigor, miaco, jsrlf, JAESO; madre, tórax, cunra, rbnjg, JEGER; hogar, lirio, mapre, nzrsf, JELUA; valor, rosál, lorna, mrbqt, JIDUO; capaz, bambú, vaípe, drnzj, JILPE; banco, ansia, úmaso, brqmt, JIUTA; corte, dardo, trida, nlrqt, JOMBE; campo, trapo, ruezo, lgcrd, JOSMO; reina, palmo, meper, sldfp, JUALE; leche, funda, nonte, mdtbj, JUCIL; papel, fresa, leluz, lrsqg, JUFÚS; clase, suela, firia, bznrd, LAFLA; coche, boina, vomón, lbctr, LÁVOX; santo, casta, mucor, nbmlf, LEDUL; joven, gasto, ruena, nzftr, LEPAZ; prisa, saldo, drolo, grjds, LIRGI; medio, secta, fluja, lsgfd, LITÉN; buena, jarra, pasno, jdgfz, LOAJO; fuego, farol, goren, gfrdm, LOPIL; grupo, verbo, badón, fnbgs, LORME; llave, andén, bejua, mgqbn, LUATE; sitio, cauce, arnia, rlfnd, LUSUR; falta, faena, fazón, dpqtm, MEIPE; noche, grifo, legor, mrbtf, MILGE; feliz, gruta, tabez, tzpgl, MIVOR; mujer, limbo, gosál, tfcrr, MOFUL; calle, aleta, maver, rbltj, MONIZ; golpe, peine, pudes, sfrbn, MOSPE; único, fobia, filza, scrzg, MOZUR; pelea, útero, bomal, rgncd, MUMET; brazo, manto, cergo, ldgrc, NAOCE; razón, cojín, broga, mrbnr, NESAX; cielo, cisne, mepla, mnrjb, NIPLO; negro, mareo, ralmo, zlngt, NOROL; reloj, cloro, runco, crrzf, NORZI; barco, pavor, orazo, rbsdl, NOVOD; color, ancla, valbo, fgrrd, NÚGUL; final, muslo, prana, tnbms, NULGE; vista, tacón, lorio, srfrnm, OBUAL; mayor, mueca, baero, qpnjf, OCOBO; sueño, ámbra, croja, cmrdt, ODUER; fácil, cuota, delle, zqbmc, OGEFA; orden, sodio, muvor, dcpng, OJUBÁ; causa, fogón, plame, drcpf, OLZIO; hielo, trama, mocer, dmqbp, ATECE; oeste, fieras, prujo, dslmn, ORSIN; trago, rigor, ficle, lfrqs, ORZUN; traje, cutis, pusto, trfbd, OSIZO; baile, vejez, mante, stldb, OTURI; señal, pudor, faubo, lfrnt, ÓVAZU; total, pompa, rargo, gzdrrn, OZABO; señor, musgo, vimal, dngtc, OZIAL; novia, tripa, carre, qrcmg, PAFIE; punto, fauna, mobun, mgcjt, PEBIO; nuevo, pasto, relío, rnlmf, PECRE; honor, fibra, zuche, qdjzp, PIEMO; banda, prosa, fánel, tqjdc, CILEA; doble,

folio, erdin, sqcfm, PIMUA; gusto, miope, mogre, fpzrj, PLIJE; lleno, momia, gopia, tpbqg, PLIME; parte, vigor, miaco, jsrlf, PÓPOZ; madre, tórax, cunra, rbnjg, PUCIZ; hogar, lirio, mapre, nzrsf, PUMUR; valor, rosál, lorna, mrbqt, PURVI; capaz, bambú, vaípe, drnzj, REBAZ; banco, ansia, úmaso, brqmt, REBUR; corte, dardo, trida, nlrqt, REMPE; campo, trapo, ruezo, lgcrd, REMÚN; reina, palmo, meper, sldfp, RIJUR; leche, funda, nonte, mdtbj, ROESO; papel, fresa, leluz, lrsqg, RUFUR; clase, suela, firia, bznrd, RUPAE; coche, boina, vomón, lbctr, RUPUO; santo, casta, mucor, nbmlf, RURME; joven, gasto, ruena, nzftr, SASIL; prisa, saldo, drolo, grjds, SAZUO; medio, secta, fluja, lsgfd, SECRI; buena, jarra, pasno, jdgfz, SEPUE; fuego, farol, goren, gfrdm, SEUNE; grupo, verbo, badón, fnbgs, SIBÚN; llave, andén, bejua, mgqbn, SICIR; sitio, cauce, arnia, rlfnd, SILUR; falta, faena, fazón, dpqtm, SOBLI; noche, grifo, legor, mrbtf, TAJUR; feliz, gruta, tabez, tzpgl, TELMÚ; mujer, limbo, gosál, tfcrr, LEFUR; calle, aleta, maver, rbltj, TIBUI; golpe, peine, pudes, sfrbn, TISCE; único, fobia, filza, scrzg, TISDO; pelea, útero, bomal, rgncd, TUGER; brazo, manto, cergo, ldgrc, TULME; razón, cojín, broga, mrbnr, TULZA; cielo, cisne, mepla, mnrjb, TURVI; negro, mareo, ralmo, zlngt, TUSGA; reloj, cloro, runco, crrzf, UCOBO; barco, pavor, orazo, rbsdl, UDAFE; color, ancla, valbo, fgrrd, UFAVI; final, muslo, prana, tnbms, UFUTA; vista, tacón, lorio, srfrnm, UGABO; mayor, mueca, baero, qpnjf, ÚGEHO; sueño, ámbra, croja, cmrdt, ULFIR; fácil, cuota, delle, zqbmc, ULGER; orden, sodio, muvor, dcpng, ULINA; causa, fogón, plame, drcpf, ÚMIVO; hielo, trama, mocer, dmqbp, UPACI; oeste, fieras, prujo, dslmn, URRIZ; trago, rigor, ficle, lfrqs, UTRÁN; traje, cutis, pusto, trfbd, UZIJO; baile, vejez, mante, stldb, VADAZ; señal, pudor, faubo, lfrnt, VÁVIL; total, pompa, rargo, gzdrrn, VÍDOY; señor, musgo, vimal, dngtc, VIRMI; novia, tripa, carre, qrcmg, VOFUR; punto, fauna, mobun, mgcjt, VOHÓN; nuevo, pasto, relío, mlfm, VORRI; honor, fibra, zuche, qdjzp, VOSÍA; banda, prosa, fánel, tqjdc, VUCEL; doble, folio, erdin, sqcfm, VUGEO; gusto, miope, mogre, fpzrj, VUJÓN; lleno, momia, gopia, tpbqg, VUOBA; parte, vigor, miaco, jsrlf, VUOJA; madre, tórax, cunra, rbnjg, TLUPO; hogar, lirio, mapre, nzrsf, XACOA; valor, rosál, lorna, mrbqt, YORPO; capaz, bambú, vaípe, drnzj, ZACUO; banco, ansia, úmaso, brqmt, ZADAO; corte, dardo, trida, nlrqt, ZIECA; campo, trapo, ruezo, lgcrd, ZIOVE; reina, palmo, meper, sldfp, ZOLIL; leche, funda, nonte, mdtbj, ZOROL; papel, fresa, leluz, lrsqg, FENDU

Received January 16, 2018

Revision received August 20, 2018

Accepted August 21, 2018 ■

Unveiling the boost in the sandwich priming technique



Quarterly Journal of Experimental Psychology
1–12
© Experimental Psychology Society 2021
Article reuse guidelines:
sagepub.com/journals-permissions
DOI: 10.1177/17470218211055097
qjep.sagepub.com



María Fernández-López¹, Colin J Davis² , Manuel Perea^{1,3} ,
Ana Marcet¹ and Pablo Gómez⁴

Abstract

The masked priming technique (which compares #####-house-HOUSE vs. #####-fight-HOUSE) is the gold-standard tool to examine the initial moments of word processing. Lupker and Davis showed that adding a pre-prime identical to the target produced greater priming effects in the sandwich technique (which compares #####-HOUSE-house-HOUSE vs. #####-HOUSE-fight-HOUSE). While there is consensus that the sandwich technique magnifies the size of priming effects relative to the standard procedure, the mechanisms underlying this boost are not well understood (i.e., does it reflect quantitative or qualitative changes?). To fully characterise the sandwich technique, we compared the sandwich and standard techniques by examining the response times (RTs) and their distributional features (delta plots; conditional-accuracy functions), comparing identity versus unrelated primes. The results showed that the locus of the boost in the sandwich technique was two-fold: faster responses in the identity condition (via a shift in the RT distributions) and slower responses in the unrelated condition. We discuss the theoretical and methodological implications of these findings.

Keywords

Masked priming; lexical decision; visual word recognition; interactive activation

Received: 15 June 2021; revised: 16 September 2021; accepted: 19 September 2021

Forster and Davis's (1984) masked priming technique has become the gold-standard tool for researchers interested in the initial moments of letter and word processing. This technique allows researchers to manipulate the relationship between a forwardly masked and briefly presented prime (approximately 30–50 ms) and a target item, while minimising the non-automatic processes induced by visible primes (see Grainger, 2008, for a review). The type of prime–target relationship depends on the ingenuity of the researchers' working hypothesis (e.g., to cite three instances: phonological priming: #####-hint-MINT vs. #####-pint-MINT; translation priming: #####-maison-HOUSE vs. #####-appel-HOUSE; and identity priming: #####-house-HOUSE vs. #####-fight-HOUSE). Proof of the ubiquity of this technique is that it has been employed in a wide range of approaches, from behavioural research—with response times (RTs) and accuracy as dependent variables (e.g., Forster et al., 1987; Grainger & Ferrand, 1994)—to computational modelling (e.g., Gomez et al., 2013) and cognitive

neuroscience (e.g., functional magnetic resonance imaging [fMRI]: Dehaene et al., 2001; enterprise resource planning [ERPs]: Grainger & Holcomb, 2009; magnetoencephalography [MEG]: Monahan et al., 2008).

Because of the short stimulus-onset asynchrony between prime and target, the magnitude of masked priming effects is small, in the range of 10–50 ms in behavioural studies. Thus, a problem faced by researchers using this technique is that it may not be sensitive enough to detect subtle manipulations or interactions (e.g., *light*-*LIGHT* vs. *liht*-*LIGHT*; see Grainger, 2008). This issue

¹Universitat de València, València, Spain

²University of Bristol, Bristol, UK

³Universidad Nebrija, Madrid, Spain

⁴California State University, San Bernardino, San Bernardino, CA, USA

Corresponding author:

Pablo Gómez, California State University, San Bernardino, Palm Desert, 5500 University Parkway, San Bernardino, CA 92407, USA.
Email: Pablo.Gomez@csusb.edu

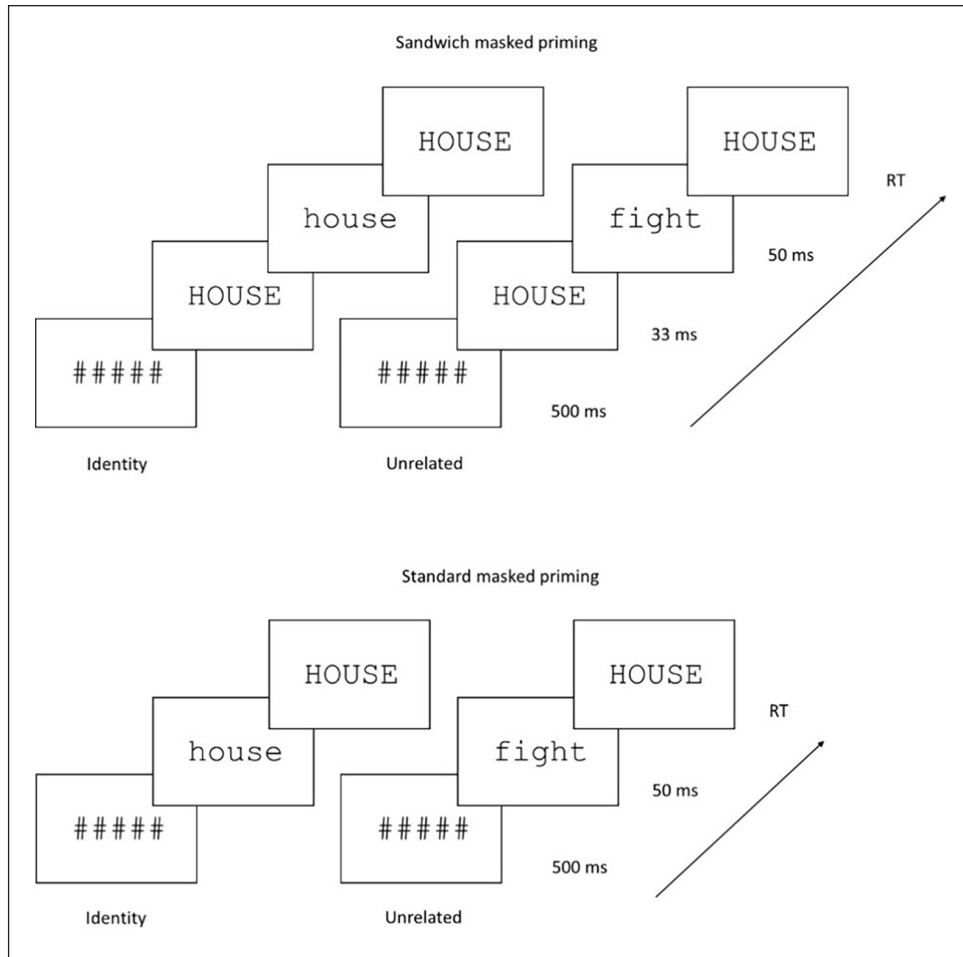


Figure 1. Sequence of events with the sandwich masked priming technique (top panel) and the standard masked priming technique (bottom panel).

is more prominent when designing experiments with developing readers or special populations (e.g., readers with dyslexia or deaf readers), for whom there tends to be more within-participant variability in the RTs compared with university students (i.e., the typical participants recruited in research studies; see L  t   & Fayol, 2013). One approach to overcome this intrinsic limitation of masked priming is statistical, by conducting over-powered large-scale experiments (e.g., Adelman et al.'s, 2014, mega-study with more than 1,000 participants, each responding to 840 stimuli). While highly valuable, this approach is costly in resources and in time, and it does not fully address the within-participant variability.

Another approach is to modify the technique to magnify the effects; this was a basic idea underlying the inception of the Lupker and Davis (2009) sandwich technique. The sandwich technique consists of a modification of the standard masked priming paradigm. The forward mask is followed by an initial masked prime identical to the target (the so-called pre-prime), followed by the prime of interest (see Figure 1 for an example). As in the standard

technique, the target replaces the prime and is displayed until the participant responds. Note that both methods involve a prime that is pre-masked (by the pre-prime in the sandwich method and by the hash marks in the standard method) and post-masked (by the target in both methods), so strictly speaking, both procedures are forms of masked priming; hence, in the remainder of the article, we will refer to them as *standard* and *sandwich* techniques.

As first shown in a series of form-priming experiments conducted by Lupker and Davis (2009), the sandwich technique produces greater priming effects than the standard technique. This boost has since been replicated with various types of prime–target relationships (e.g., see Comesa  a et al., 2016; Perea et al., 2014; Stinchcombe et al., 2012; Trifonova & Adelman, 2018, for behavioural evidence, and Ktori et al., 2012, for electrophysiological evidence). Indeed, because of the larger priming effects, the sandwich methodology has become the preferred option in many recent masked priming studies (e.g., Campos et al., 2021; Colombo et al., 2020; Gutierrez-Sigut et al., 2017; Ktori et al., 2014; Lupker et al., 2015, 2020a,

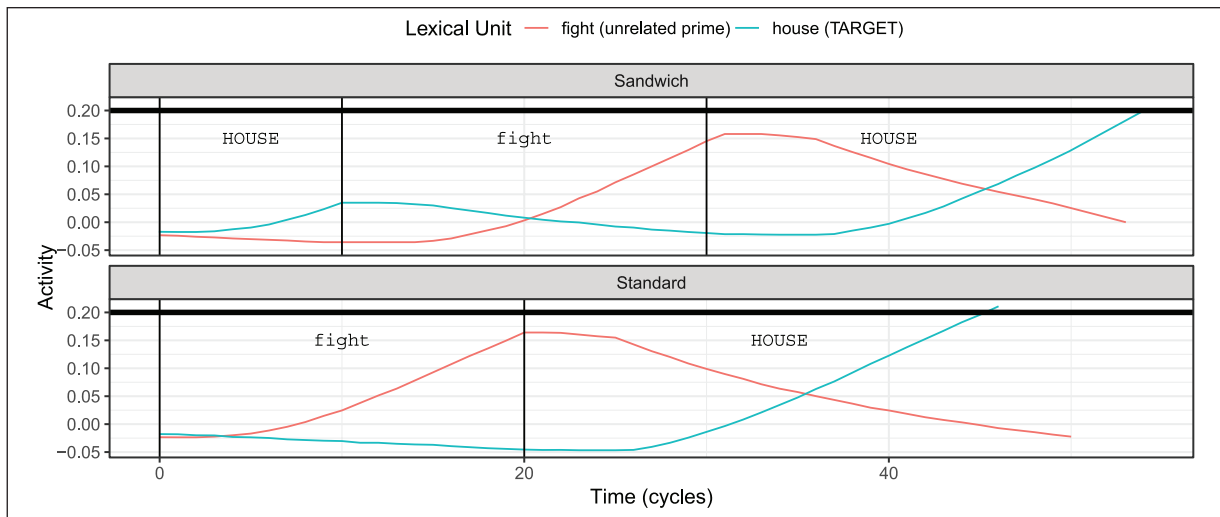


Figure 2. Example of the initial processing cycles of an unrelated trial in the sandwich technique (top panel, pre-prime: house [10 cycles], prime: fight [20 cycles], and target [until response]) and in the standard masked priming technique (bottom panel, prime: fight [20 cycles] and target [until response]) in a leading competitive network model of visual word recognition (SCM; C. J. Davis, 2010). The thick horizontal line represent a response threshold, and the coloured lines represent the activation levels of the corresponding lexical units.

2020b; Meade et al., 2020; Taikh & Lupker, 2020; Ziegler et al., 2014).

Notwithstanding its utility in amplifying priming effects, the sandwich technique requires some further examination. The specific mechanisms at play and how they differ from the standard technique, have not been fully elucidated (see Lupker & Davis, 2009 vs. Trifonova & Adelman, 2018). Thus, it is not yet well understood whether the priming effects in the sandwich technique are simply magnified or, instead, the underlying processes are fundamentally altered. Exploring this question is the main goal of the present work. In the original Lupker and Davis (2009) experiments, task procedure (sandwich vs. standard) was manipulated between subjects, making it difficult to compare the nuances of the two techniques. Trifonova and Adelman (2018) have been the only researchers to date who employed a within-subject design for task procedure. However, their form-priming experiments focused on whether the characteristics of the pre-prime influenced word processing (e.g., how the pre-prime *PROTECT* influences the processing of the target *PROJECT*) rather than a direct comparison of the two techniques using their most common set-up as in Figure 1.

To compare the two techniques, we manipulated the procedure (standard vs. sandwich) in a within-subject design where we compared the largest and most robust form of masked priming: identity priming (i.e., identity vs. unrelated primes). Our comparison aims to shed light on the processing differences between the two techniques. To do so, we also examine the distributional features of the RT data and the relative speed of correct and error responses in the two paradigms.

For identity primes, the activation from the pre-primes in the sandwich technique would naturally help word processing: the target word is presented as a pre-prime and as a prime (e.g., compare pre-prime: *HOUSE*, prime: *house*, target: *HOUSE* in the sandwich method vs. prime: *house*; target: *HOUSE* in the standard method). Thus, word recognition times in the identity condition should be faster in the sandwich technique than in the standard technique. For the unrelated primes, the presence of the pre-prime in the sandwich technique does not generate such a straightforward prediction (e.g., pre-prime: *HOUSE*, prime: *fight*, target: *HOUSE* in the sandwich technique vs. prime: *fight*, target: *HOUSE* in the standard technique). As suggested by Trifonova and Adelman (2018), in the standard technique one would expect unrelated primes to push the target's activation below resting level after prime offset in interactive activation models of word recognition (e.g., SCM, C. J. Davis, 2010; see bottom panel of Figure 2). In contrast, because of the residual activation caused by the pre-prime in the sandwich technique, the target's activation would be higher than with the standard technique after prime offset (see top panel of Figure 2). Thus, one would predict that unrelated pairs would be identified faster in the sandwich than in the standard technique—note that this mechanism would reduce the boost (i.e., reduce the priming effect) of the sandwich technique.

The above pattern was confirmed by simulations on the spatial coding model (C. J. Davis, 2010). We used the target words of C. J. Davis and Lupker's (2016) Experiment 2—they were all five-letter words—paired with identity versus unrelated primes.¹ For the sandwich method, we simulated a scenario with a 33-ms pre-prime followed by a

50-ms prime, whereas for the standard technique, we simulated a scenario with a 50-ms prime. We found a greater identity priming effect in the sandwich technique than in the standard technique (69 vs. 47 processing cycles). This boost was exclusively due to the identity condition (an advantage of 24 processing cycles in the sandwich technique). For the unrelated condition, the sandwich technique produced slightly faster responses than the standard technique (an advantage of two processing cycles in the sandwich technique). Of note, in a form priming experiment, Trifonova and Adelman (2018, Experiment 2) found the opposite trend for unrelated pairs in the error rates (i.e., better performance in the standard technique); if this advantage for the standard technique becomes the dominant pattern of results, this would suggest that the way models simulate masked priming with the standard and sandwich techniques many need to be reevaluated.

To get a richer picture of the differences between the two techniques, we also aim to fully describe the distributional features of the RTs in the sandwich technique and standard methods; in the present work, we report delta plots (De Jong et al., 1994) and conditional accuracy functions (Bonnet & Dresp, 1993; Ollman, 1977). Critically, the study of RT distributions offers rich information to elucidate the time course of an effect (Balota et al., 2008; see also Heathcote et al., 1991; Ratcliff, 1979; Townsend & Ashby, 1983, for early research), because many theories of word processing make claims of the sequence and timing of mental operations, as can be seen in Figure 2. However, all previous studies with the sandwich technique only focused on the analysis of the mean RTs.

In the standard technique, masked identity priming has repeatedly shown a shift between the RT distributions of the identity and unrelated conditions (i.e., identity priming), with the magnitude of identity priming being close to that of the prime duration (e.g., Gomez et al., 2013; Gomez & Perea, 2020; Perea et al., 2018, for a similar finding). This pattern supports Forster's (1998) claim that the identity prime produces a "head start" advantage when processing the target word: The identity pair house-HOUSE would enjoy an advantage over the unrelated pair fight-HOUSE of approximately the prime duration. In contrast, when using visible primes (150-ms prime duration), identity priming elicits a shift in the RT distributions and a change in their shape as it is more skewed for the unrelated condition. This pattern is consistent with the idea that unmasked visible primes affect core decision processes (see Gomez et al., 2013). Here, we tested whether masked identity priming effects in the sandwich technique also reflect a shift in the RT distributions (i.e., a head-start advantage, as in the standard masked priming technique, simply of a larger size) or whether they are larger in the higher quantiles (i.e., activation of the lexical units grow over time as a consequence of the identity primes). The former outcome would support the view that the two

techniques tap the same core processes, whereas the latter result would imply that the sandwich technique taps core processes other than a "head-start" advantage.

In sum, to provide the first full description of the sandwich technique, we designed a lexical decision experiment that examined the similarities and differences in masked identity priming with the standard and sandwich techniques. The findings of this study are also relevant to inform us of whether the sandwich technique taps similar encoding processes (i.e., a head-start advantage) as the standard priming technique or whether the two tasks are more different than originally thought. Keep in mind that, in the sandwich technique, the pre-prime presentation adds complexity to the information processing; that is, the system must handle with three strings in brief succession (see Forster, 2009; Trifonova & Adelman, 2018). To rephrase our research question, an apt analogy might be that the two methods can be thought of as measuring devices for cognitive processes and the present study aims to understand whether these devices are qualitatively different as they detect different signals, or quantitatively different, as they may magnify the same signal to different levels.

Method

Participants

A total of 36 undergraduate students from the University of Valencia volunteered to participate in the experiment—this ensured 2,160 observations per condition, in line with Brysbaert and Stevens's (2008) suggestion for masked priming experiments. All participants were native speakers of Spanish with normal/corrected vision and no history of reading disorders. This study was approved by the Experimental Research Ethics Committee of the University of Valencia, and all participants signed an informed consent form.

Materials

We selected 240 word targets of 5 letters, which were extracted from Fernández-López et al.'s (2019) Experiment 1. The average Zipf frequency was 4.15 (range 3.74–4.61) and the mean OLD20 was 1.41 (range: 1.00–2.65) in the EsPal Spanish database (Duchon et al., 2013). Each target word (e.g., PASTA [pasta]) was paired with an identity prime (e.g., pasta) or an unrelated prime (e.g., suelo [floor]). The pre-primes in the sandwich method were always the same as the targets (see Figure 1). To act as foils, we selected 240 orthographically legal pseudowords. These were also extracted from the Fernández-López et al.'s (2019) Experiment 1. Each target nonword (e.g., JARSA) was paired with an identity prime (jarsa) or an unrelated prime (ruezo). We created four lists to counterbalance the prime–target combinations across the two

Table 1. Mean correct response times (in ms) and accuracy (in brackets) for words and pseudowords with the standard and sandwich masked priming techniques.

	Standard technique		Sandwich technique	
	Identity	Unrelated	Identity	Unrelated
Words	589 (96.9)	629 (93.9)	563 (96.9)	642 (91.2)
Pseudowords	704 (91.6)	717 (92.6)	695 (92.7)	709 (92.0)

priming conditions and the two techniques—they are available at the OSF link <https://osf.io/83x9t>

Procedure

Testing took place in groups of 4–6 participants in a quiet laboratory room. Computers equipped with DMDX software (Forster & Forster, 2003) displayed the stimuli and registered the timing/accuracy of the responses. Each trial started with a pattern mask (#####) displayed for 500 ms in the centre of a computer screen. In those trials with the sandwich technique, the uppercase target stimulus was presented for 33 ms as a pre-prime, followed by a 50-ms lowercase prime stimulus which, in turn, was replaced by the uppercase target stimulus—this was presented until response or 2 s had elapsed. In the trials with the standard technique, the pattern mask was followed by the 50-ms prime, and then replaced by the target (see Figure 1). All stimuli were presented in 14-pt Courier New black font on a white background. Participants were instructed to decide whether the uppercase stimulus was a word or not by pressing a green button (“word”) or a red button (“non-word”) with their index fingers. Both speed and accuracy were stressed in the instructions. Instructions did not mention the existence of any briefly presented primes. A total of 16 practice trials preceded the 480 experimental trials. Each participant received a random sequence of trials. The session lasted 18–22 min.

Results

Error responses and anticipations (RTs faster than 250 ms: 2 correct responses and 4 incorrect responses) were omitted from the latency analyses. Timeouts after 2.5 s were categorised as errors (there were 32 timed out responses and 1,091 error responses out of 17,280 trials). Table 1 displays the mean RTs for correct responses and the accuracy in each condition. We first present the inferential analyses using Bayesian linear mixed-effects models, and then we present the exploratory data analyses.

Bayesian linear mixed-effects analyses

The RT and accuracy data were modelled with Bayesian linear mixed-effects models in R (R Core Team, 2020)

using the *brms* package (Bürkner, 2020). This package, which uses the programming language Stan as its core, allows us to fit the maximal random effect structure models allowed by the experimental design (see Barr et al., 2013). The fixed effects were relation (identity vs. unrelated; coded as -0.5 and 0.5) and procedure (sandwich vs. standard; coded as -0.5 and 0.5). As usual in masked priming experiments, we conducted separate analyses for word and nonword trials.

The fitted model was as follows:

```
Dep. variable [RT or accuracy] ~ relation * procedure + (1 + relation * procedure | subject) + (1 + relation * procedure | item)
```

We used the ex-Gaussian family function to model the RT data and the Bernoulli family function ($1 = \text{correct}$, $0 = \text{incorrect}$) to model the accuracy data. For the fits of each model we employed four chains, each with 5,000 iterations (1,000 warmup + 4,000 sampling). All R values were 1.00, thus meaning that the four chains converged successfully. In the output, the Bayesian linear mixed-effects models indicate the estimate of each fixed effect, its standard error, and its 95% credible interval. For inference purposes, those intervals that did not include zero were interpreted as significant. In the case of a significant interaction, we computed the simple effect tests with the *emmeans* package (Lenth et al., 2020)—this package provides the 95% highest (posterior) density (HPD) interval for each effect. Note that this method yields what is called Bayesian shrinkage, which means that posterior estimates might be shifted from the empirical effects due to the prior and the data from other conditions.

Word targets

Response times were faster for identity pairs than for unrelated pairs ($b = 72.35$, $SE = 3.78$, 95% credible interval [CrI] [64.58, 79.56]) and responses were faster in the sandwich than in the standard technique ($b = 25.65$, $SE = 2.70$, 95% CrI [20.74, 31.29]). As expected, there was clear evidence of an interaction between the two factors ($b = -32.76$, $SE = 4.09$, 95% CrI [-40.80, -24.78]). This interaction reflected that, for identity pairs, responses were substantially faster in the sandwich than in the standard technique (95% HPD [-31.29, -20.7]); in contrast, for unrelated pairs, responses were faster in the standard than in the sandwich technique (95% HPD [0.75, 13.6]).

The analyses on the accuracy data showed that participants were more accurate for identity pairs than for unrelated pairs ($b = -1.16$, $SE = 0.19$, 95%CrI [-1.52, -0.79]). While the 95% credible intervals of the other parameters included zero, accuracy was lower in the sandwich than in the standard technique ($b = 0.18$, $SE = 0.23$, 95%CrI [-0.28,

0.65]) and the size of the identity priming effect was greater in the sandwich than in the standard technique (interaction: $b=0.40$, $SE=0.26$, 95%CrI [-0.11, 0.90]).

Nonword targets

Responses to nonwords were faster in the trials with the sandwich technique than in the trials with the standard technique ($b=17.54$, $SE=3.67$, 95% CrI [10.37, 24.72]). We also found a small advantage of identity over unrelated pairs ($b=9.83$, $SE=5.07$, 95% CrI [0.22, 19.93]). There were no clear signs of an interaction between the two factors ($b=-4.21$, $SE=5.11$, 95% CrI [-14.30, 5.7]).

The analyses on the accuracy data did not show any effects (relation: $b=-0.14$, $SE=0.16$, 95% CrI [-0.45, 0.18]; procedure: $b=-0.14$, $SE=0.16$, 95% CrI [-0.45, 0.17]; relation*procedure: $b=0.21$, $SE=0.21$, 95% CrI [-0.20, 0.61]).

To summarise, the results with the Bayesian linear mixed-effects models on word trials showed that, for identity primes, the presentation of the target word in the sandwich technique yielded faster responses than the standard technique (HOUSE-house-HOUSE < house-HOUSE). In contrast, the sandwich technique slowed down word responses for unrelated primes (fight-HOUSE < HOUSE-fight-HOUSE). Thus, the boost in the sandwich technique is caused by faster responses in the identity condition and slower responses in the unrelated condition. Regarding the nonword trials, we found faster responses in the sandwich technique, accompanied by a small identity priming effect.

Exploratory data analyses

To further examine the nature of the priming effects underlying the sandwich and standard techniques, we compared the empirical RT distributions for both techniques. We focused on delta plots and conditional accuracy functions, as they describe how the latency distributions differ. Note that these methods do not have established inferential properties, so in our discussion we focus only on large effects.

Delta plots. Delta plots are frequently used to display how a latency effect (e.g., identity priming, or task) evolves across time (see De Jong et al., 1994; Ridderinkhof et al., 2004). These plots are constructed as follows:

1. We obtain the RTs at the desired quantiles (here we use .1, .2, .3, .4, .5,9) for every participant for each of the conditions to be compared, only for correct responses.
2. We average the RTs obtained in Step 1 across participants (these averaged quantiles are also called vincentiles in the RT literature; see Ratcliff, 1979).

3. For each of the quantiles in the vincentiles (a) we find the average between the two conditions, and (b) then compute the differences (i.e., the delta) preserving the sign.
4. The last step is to plot the points (as many points as there are quantiles), with the averages (Step 3a) on the x -axis, and the delta (Step 3b) on the y -axis.

As is evident from Step 3b, delta plots are based on residual quantiles (i.e., RTs at quantiles in condition A minus condition B), and hence can provide us with some indication of the temporal dynamics of an effect when interpreted within the lens of process models. For example, a flat line at around $y=50$ ms would mean that there is a 50-ms shift in the RT distributions—this could be interpreted as faster encoding times in evidence accumulation models (see Gomez et al., 2013, for an example of such interpretation in the standard masked priming technique). Alternatively, an ascending function would indicate that the effect grows for slower responses, and it could be interpreted as a difference in the rate of evidence accumulation.

As just indicated, in delta plots, one needs to contrast two conditions; in our analysis of the two priming methods, it makes sense to perform two separate comparisons that are presented in two different delta plots: (a) the effect of prime type in the standard and sandwich techniques (across condition comparison: unrelated vs. identity; Figure 3), and (b) the effect of procedure for identity and unrelated pairs (within-condition comparison: standard vs. sandwich; Figure 4).

Effect of prime type (Figure 3). For word trials, when calculating priming effects as RT for unrelated primes minus RT for identity primes, we found a relatively flat line for the identity priming effect in the standard technique (all points lie between 38 and 51 ms), thus replicating earlier research. In contrast, in the sandwich technique, the identity priming effect was not only greater than in the standard technique (as attested by the higher values in the y -axis), but it progressively increased from the .1 quantile (63 ms) to the .7 quantile (95 ms)—it decreased slightly in the .9 quantile, though. For the nonword trials, we found a small identity priming effect that slowly grew across quantiles, with a steeper slope in the sandwich technique.

Effect of technique (Figure 4). For word trials, when calculating the effect of technique as *standard* minus *sandwich*, we found an approximate flat line for the identity primes (around 25 ± 5 ms for all quantiles). This pattern indicates that responses were overall *faster* in the sandwich than in the standard technique. In contrast, we found a line that lies on the negative numbers for the unrelated primes (i.e., responses were *slower* in the sandwich than in the standard technique)—note that

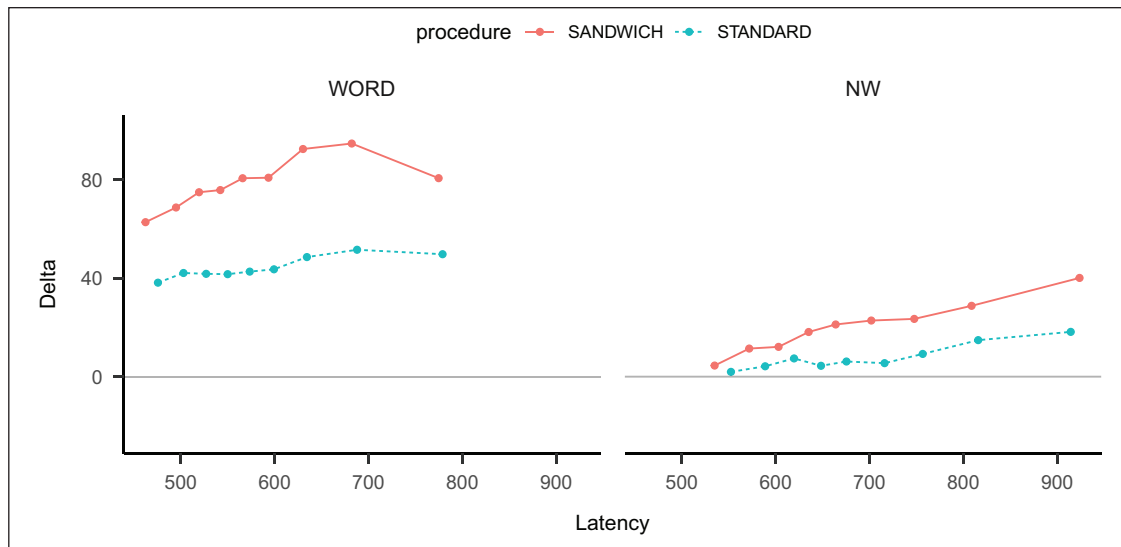


Figure 3. Delta plot comparing unrelated primes to identity primes. The points represent the difference $RT_{unrelated} - RT_{identity}$ at the .1, .2, .3, .4, .5, .6, .7, .8, and .9 quantiles.

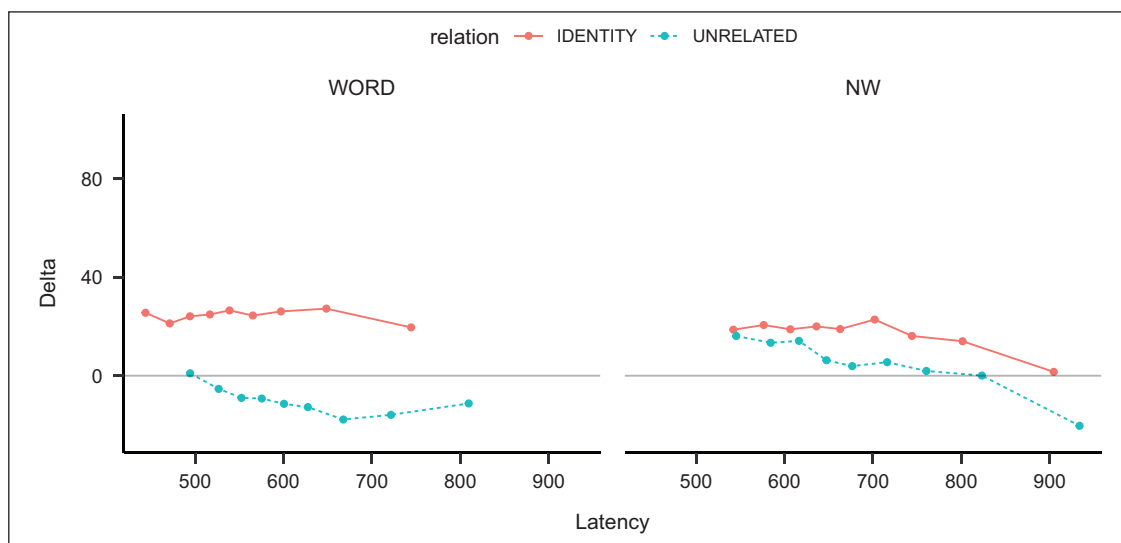


Figure 4. Delta plot comparing the standard and sandwich techniques. The points represent the residual of $RT_{standard} - RT_{sandwich}$ at the .1, .2, .3, .4, .5, .6, .7, .8, and .9 quantiles.

this difference effect grew as quantiles increased from .1 quantile (−.5 ms) to .8 quantile (−.16 ms; i.e., a negative slope). We found functions with negative slopes for the nonword targets, but with points above zero for the most part. This pattern reflected that responses were faster in the sandwich technique, but this advantage decreases (or even reversed, in the unrelated condition) for the slower responses.

Note that delta plots highlight the dynamics of effects across different points of the latency distributions. However, they do not convey information about accuracy as they are only built with correct responses; hence, we

resort to conditional accuracy functions to explore accuracy issues as explained below.

Conditional accuracy functions. Conditional accuracy functions (Bonnet & Dresch, 1993; Ollman, 1977) allow us to visualise how accuracy evolves as latencies increase. These plots were generated as follows:

1. For each type of item and each participant, we divide all responses into 5 equal-sized bins based on their latency; to do so, we included both correct and error responses.

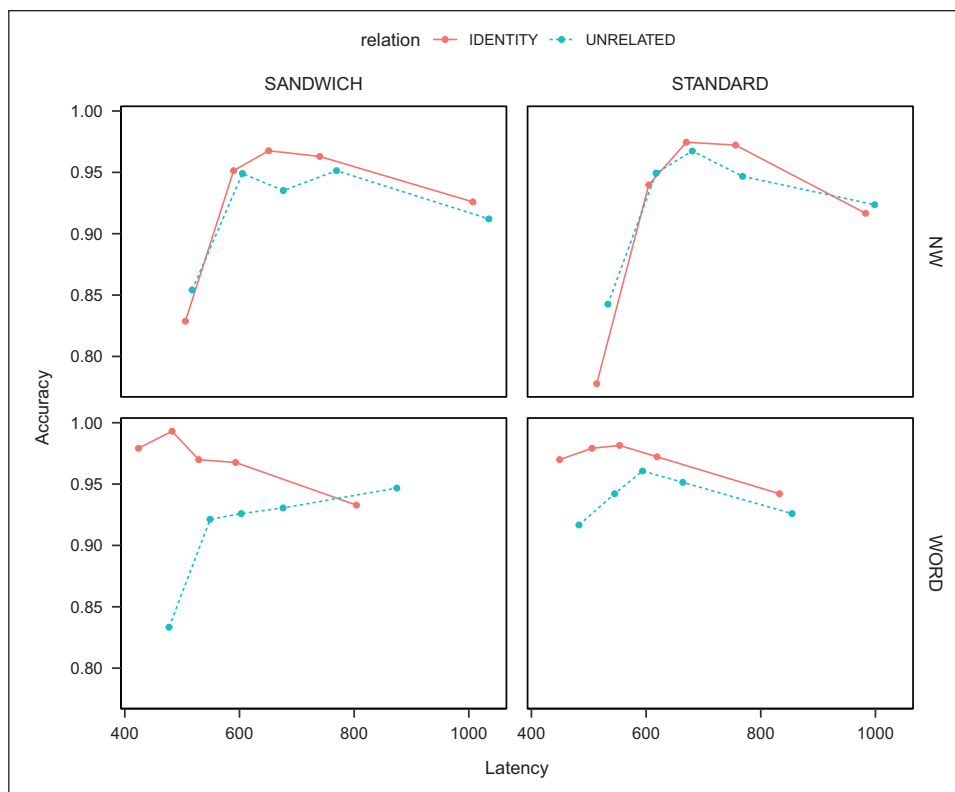


Figure 5. Conditional accuracy functions, the points represent the accuracy and average RT of the responses within equal-sized bins (20% of responses per bin).

2. We averaged the RTs and accuracies for each bin across all participants, and then calculate (a) the average RT for each of the bins and (b) the accuracy within each bin.
3. The last step is to plot the points (as many points as there are quantiles), with the averages (Step 3a) on the x-axis, and the conditional accuracies (Step 3b) on the y-axis.

As is evident from Step 3b, conditional accuracy functions can indicate whether the accuracy in the responses for a given condition varies across the latency of the responses to such condition. For example, a flat line at around $y = .90$ would mean that the accuracy for that condition (90%) is equal for fast and for slow responses. In addition, an ascending function would indicate that the fast responses tend to be less accurate, and a descending function would indicate that slow responses tend to be less accurate.

The panels in Figure 5 show the conditional accuracy for all the conditions in our experiment. Two patterns are worth highlighting: for words, the fastest 20% of responses in the unrelated condition in the sandwich method tend to be less accurate than slower responses, which is reversed for the identity condition. For nonwords, the conditional accuracy functions are remarkably similar in the two

techniques: fastest responses tend to be less accurate for the identity primes than for the unrelated primes, but this pattern reverses for the slower responses. This is consistent with the idea that familiarity influences lexical decision (e.g., Masson & Bodner; cf compound-cues), particularly for fast responses, which could be sensitive to the orthographic match of the prime and target. A match could bias for “word” responses, and fast responses might be most affected by this bias.

Discussion

An increasing number of masked priming researchers now opt for the sandwich technique due to its promise of greater effect sizes than the standard technique. However, the mechanisms underlying this boost are still not well understood. To fully characterise the sandwich technique, we conducted a lexical decision experiment comparing the sandwich and standard techniques in a within-subject design. We focused on masked identity priming effects for three reasons: (a) theoretical models offer straightforward predictions (Forster, 1998); (b) it has already been modelled with the standard masked priming technique (Gomez et al., 2013); and (c) the magnitude of the effects is much larger than with other types of prime–target relationships (i.e., we may capture more easily potential interactions).

To elucidate whether the differences between the standard and sandwich methods were merely quantitative or also qualitative, we examined the locus of the boost in the sandwich technique and the time course of the effects for both techniques via delta plots and conditional accuracy functions.

For word trials, we found the expected boost in the sandwich technique (a 79-ms priming effect with the sandwich technique vs. a 40 ms priming effect with the standard method; i.e., a 39-ms boost), thus replicating earlier research (see Lupker & Davis, 2009, for the first demonstration). Critically, the use of a within-subject design allowed us to discover that while the major component of the 39-ms boost in the magnitude of priming effects in the sandwich technique was due to the faster responses to identity pairs (26 ms: 563 vs. 589 ms in the sandwich and standard techniques, respectively), there was also a component due to the slower responses to unrelated pairs (13 ms: 642 vs. 629 ms, in the sandwich and standard techniques, respectively)—note that Trifonova and Adelman (2018) also found an effect in this direction for unrelated pairs in the accuracy data. As indicated in the Introduction, C. J. Davis's (2010) spatial coding model would have predicted not only substantially faster responding in the identity condition in the sandwich than in the standard technique, but also slightly faster responding in the unrelated condition. Thus, the effect of the unrelated priming conditions predicted by the spatial coding model does not match the observed data. We found that the responses were slower in the sandwich technique than in the standard technique after an unrelated prime. The discrepancy in the unrelated trials suggests that the pre-prime in the sandwich technique may influence target processing in a way that cannot be captured in the present implementations of masked priming in interactive activation models like the spatial coding model.

In addition, the exploratory data analyses (both delta plots and conditional accuracy functions) offer valuable information on the underlying processing differences between the standard and sandwich techniques. For word targets, the delta plot on the standard technique showed, as expected, that masked identity priming reflects essentially a shift in the RT distributions (see Figure 3; see also Gomez et al., 2013; Gomez & Perea, 2020). This pattern is entirely consistent with the commonly shared assumption that masked identity primes have a “head start” advantage during an early encoding stage (i.e., a “savings” effect; Forster, 1998; see also Gomez et al., 2013). In contrast, the delta plot with the sandwich technique showed a different picture: the magnitude of masked identity priming grows stronger in the higher quantiles (see Figure 3). At first glance, one might interpret this last pattern as being due to some involvement of decision processes for identity primes in the sandwich technique. However, this interpretation misses the critical information provided by within-condition comparisons (see Jacobs et al., 1995). The

advantage of the identity primes in the sandwich technique is approximately constant along quantiles (see left panel of Figure 4), thus reinforcing the idea that the identity primes are processed similarly in the two techniques (i.e., an encoding effect; Forster, 1998). Instead, the largest differences across tasks occur in the unrelated priming condition, where there is a costlier processing in the sandwich technique as latency increases. Of note, the delta plot in the left panel of Figure 4 also allows us to visualise that the difference between the sandwich and the standard method manifests itself in both identity and unrelated primes: the sandwich technique yields faster responses for the identity primes and slower responses for the unrelated primes.

The conditional accuracy functions shown in Figure 5 provide further clues on the differences between the two techniques. The fast responses in the sandwich method for the unrelated condition are much less precise than the fast responses for any other condition in the experiment. This pattern suggests that the pre-prime interacts with the prime in ways that affect the early moments of processing of the target. One possible explanation for the divergent pattern in the unrelated primes across technique is the following. Because of the brevity of the pre-prime, its processing may overlap with the prime, creating a jumbled orthographic code that would impair the first moments of target processing—note that this interpretation has some resemblance with the legibility effect described by C. Davis and Forster (1994). For instance, for *HOUSE-fight-HOUSE* (sandwich technique), when participants initially encounter *HOUSE-fight*, they might perceive a fused orthographic code that might initially hinder target recognition, occasionally leading to fast error responses. This early difficulty would be later overwhelmed by lexical activation.² Importantly, this is a post hoc explanation that can be interpreted in different ways. What do we mean by a “jumbled orthographic code”? Two implementations of this idea come to mind: the first one is that the orthographic features of both the pre-prime and the target are partially available, and hence there might be at least some benefit for the sandwich technique over the standard technique; unfortunately, this idea does not seem to be supported by the data because RTs for the sandwich technique are slightly longer than for the standard technique. Another way to think about this “jumbling” is that the orthographic codes mutually inhibit each, thus diminishing lexical activation for both entries. This question, however, is well beyond the scope of the present article. Nevertheless, it raises interesting issues around how information from the strings integrates versus contrast in the type of fast presentation displays in this type of paradigms. These questions will be central in future analyses of the sandwich technique, which was developed to explore rather fine-grained orthographic processes that yield much smaller effects than identity versus unrelated primes.

In sum, the present experiment has both theoretical and practical implications. On the theoretical side, the shift in

the RT distributions for identity pairs in the sandwich and standard technique strongly suggests that the extra time from the pre-prime does not alter the underlying processes beyond encoding. Notably, the story is different from unrelated primes: the pre-prime in the sandwich technique hinders the processing of unrelated prime-target pairs. Further research with a better temporal resolution (i.e., ERPs) is necessary to uncover these differences—note that the Ktori et al. (2012) ERP sandwich experiment did not include a block with the standard technique. On the applied side, one basic question is shall we continue using the sandwich technique? The answer from the present experiment is “yes.” Even in the largest manipulation in priming experiments (i.e., identity primes), we found a similar pattern in the two techniques for identity pairs—except for the expected processing advantage in the sandwich technique due to the pre-prime. At the same time, researchers should be cautious in interpreting some of the across-condition priming effects in subtle manipulations in the sandwich technique, as these effects may derive not from facilitation in the related condition but rather from some inhibition in the unrelated condition and might consider using both the identity and unrelated conditions as controls when exploring more subtle manipulation (e.g., does corn prime ACORN?). Perhaps an analogy from the physical sciences is apt. We explored whether the sandwich technique is a more powerful microscope than the standard masked priming technique. Our results indicate that a better way to think about it is as a different catalyst that triggers related but slightly processes.

Declaration of conflicting interests

The author(s) declared no potential conflicts of interest with respect to the research, authorship, and/or publication of this article.

Funding

The author(s) disclosed receipt of the following financial support for the research, authorship, and/or publication of this article: The research reported in this article has been partially supported by Grants PRE2018-083922 (M.F.-L.) and PSI2017-86210-P (M.P.) funded by MCIN/AEI/ 10.13039/501100011033, and the Grant GV/2020/074 (A.M.) from the Valencian Government.

Ethical approval

The procedures involving human participants in this study were approved by the Experimental Research Ethics Committee of the Universitat de València and they were in accordance with the Declaration of Helsinki. All participants provided written informed consent before starting the experimental session.

ORCID iDs

Colin J Davis  <https://orcid.org/0000-0001-5814-5104>

Manuel Perea  <https://orcid.org/0000-0002-3291-1365>

Ana Marcet  <https://orcid.org/0000-0001-8755-5903>

Pablo Gómez  <https://orcid.org/0000-0003-4180-1560>

Data accessibility statement



The data and materials from the present experiment are publicly available at the Open Science Framework website: <https://osf.io/83x9t>

Notes

1. The simulator is available at <http://www.pc.rhul.ac.uk/staff/c.davis/SpatialCodingModel/>
2. While the main focus of this work is the words, no complete account of the tasks can exclude a description of the responses to pseudowords. In our data, the conditional accuracy function and the delta plots are remarkably similar, with effects augmented in the sandwich condition.

References

- Adelman, J. S., Johnson, R. L., McCormick, S. F., McKague, M., Kinoshita, S., Bowers, J. S., . . . Davis, C. J. (2014). A behavioral database for masked form priming. *Behavior Research Methods*, 46, 1052–1067. <https://doi.org/10.3758/s13428-013-0442-y>
- Balota, D. A., Yap, M. J., Cortese, M. J., & Watson, J. M. (2008). Beyond mean response latency: Response time distributional analyses of semantic priming. *Journal of Memory and Language*, 59, 495–523. <https://doi.org/10.1016/j.jml.2007.10.004>
- Barr, D. J., Levy, R., Scheepers, C., & Tily, H. J. (2013). Random effects structure for confirmatory hypothesis testing: Keep it maximal. *Journal of Memory and Language*, 68, 255–278. <https://doi.org/10.1016/j.jml.2012.11.001>
- Bonnet, C., & Dresch, B. (1993). A fast procedure for studying conditional accuracy functions. *Behavior Research Methods, Instruments, & Computers*, 25, 2–8. <https://doi.org/10.3758/BF03204443>
- Brysbaert, M., & Stevens, M. (2018). Power analysis and effect size in mixed effects models: A tutorial. *Journal of Cognition*, 1(9), 1–20. <http://doi.org/10.5334/joc.10>
- Bürkner, P. C. (2020). Analysing standard progressive matrices (SPM-LS) with Bayesian item response models. *Journal of Intelligence*, 8, 5. <https://doi.org/10.3390/jintelligence8010005>
- Campos, A. D., Oliveira, H. M., & Soares, A. P. (2021). Syllable effects in beginning and intermediate European-Portuguese readers: Evidence from a sandwich masked go/no-go lexical decision task. *Journal of Child Language*, 48(4), 699–716. <https://doi.org/10.1017/S0305000920000537>
- Colombo, L., Spinelli, G., & Lupker, S. J. (2020). The impact of consonant–vowel transpositions on masked priming effects in Italian and English. *Quarterly Journal of Experimental Psychology*, 73, 183–198. <https://doi.org/10.1177/41740720121881986677636838>
- Comesaña, M., Soares, A. P., Marcet, A., & Perea, M. (2016). On the nature of consonant/vowel differences in letter position coding: Evidence from developing and adult readers. *British Journal of Psychology*, 107, 651–674. <https://doi.org/10.1111/bjop.12179>
- Davis, C., & Forster, K. I. (1994). Masked orthographic priming: The effect of prime-target legibility. *Quarterly Journal*

- of *Experimental Psychology*, 47, 673–697. <https://doi.org/10.1080/14640749408401133>
- Davis, C. J. (2010). The spatial coding model of visual word identification. *Psychological Review*, 117, 713–758. <https://doi.org/10.1037/a0019738>
- Davis, C. J., & Lupker, S. J. (2006). Masked inhibitory priming in English: Evidence for lexical inhibition. *Journal of Experimental Psychology: Human Perception and Performance*, 32, 668–687. <https://doi.org/10.1037/0096-1523.32.3.668>
- Dehaene, S., Naccache, L., Cohen, L., Le Bihan, D., Mangin, J., Poline, J., & Riviere, D. (2001). Cerebral mechanisms of word masking and unconscious repetition priming. *Nature Neuroscience*, 4, 752–758. <https://doi.org/10.1038/89551>
- De Jong, R., Liang, C.-C., & Lauber, E. (1994). Conditional and unconditional automaticity: A dual-process model of effects of spatial stimulus-response correspondence. *Journal of Experimental Psychology: Human Perception and Performance*, 20, 731–750. <https://doi.org/10.1037/0096-1523.20.4.731>
- Duchon, A., Perea, M., Sebastián-Gallés, N., Martí, A., & Carreiras, M. (2013). EsPal: One-stop shopping for Spanish word properties. *Behavior Research Methods*, 45, 1246–1258. <https://doi.org/10.3758/s13428-013-0326-1>
- Fernández-López, M., Marcet, A., & Perea, M. (2019). Can response congruency effects be obtained in masked priming lexical decision? *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 45, 1683–1702. <https://doi.org/10.1037/xlm0000666>
- Forster, K. I. (1998). The pros and cons of masked priming. *Journal of Psycholinguistic Research*, 27, 203–233. <https://doi.org/10.1023/A:1023202116609>
- Forster, K. I. (2009). The intervenor effect in masked priming: How does masked priming survive across an intervening word? *Journal of Memory and Language*, 60, 36–49. <https://doi.org/10.1016/j.jml.2008.08.006>
- Forster, K. I., & Davis, C. (1984). Repetition priming and frequency attenuation in lexical access. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 10, 680–698. <https://doi.org/10.1037/0278-7393.10.4.680>
- Forster, K. I., Davis, C., Schoknecht, C., & Carter, R. (1987). Masked priming with graphemically related forms: Repetition or partial activation? *Quarterly Journal of Experimental Psychology*, 39, 211–251. <https://doi.org/10.1080/14640748708401785>
- Forster, K. I., & Forster, J. C. (2003). DMDX: A Windows display program with millisecond accuracy. *Behavior Research Methods*, 35, 116–124. <https://doi.org/10.3758/BF03195503>
- Gomez, P., & Perea, M. (2020). Masked identity priming reflects an encoding advantage in developing readers. *Journal of Experimental Child Psychology*, 199, 104911. <https://doi.org/10.1016/j.jecp.2020.104911>
- Gomez, P., Perea, M., & Ratcliff, R. (2013). A diffusion model account of masked versus unmasked priming: Are they qualitatively different? *Journal of Experimental Psychology: Human Perception and Performance*, 39, 1731–1740. <https://doi.org/10.1037/a0032333>
- Grainger, J. (2008). Cracking the orthographic code: An introduction. *Language and Cognitive Processes*, 23, 1–35. <https://doi.org/10.1080/01690960701578013>
- Grainger, J., & Ferrand, L. (1994). Phonology and orthography in visual word recognition: Effects of masked homophone primes. *Journal of Memory and Language*, 33, 218–233. <https://doi.org/10.1006/jmla.1994.1011>
- Grainger, J., & Holcomb, P. J. (2009). Watching the word go by: On the time-course of component processes in visual word recognition. *Language and Linguistics Compass*, 3, 128–156. <https://doi.org/10.1111/j.1749-818X.2008.00121.x>
- Gutierrez-Sigut, E., Vergara-Martínez, M., & Perea, M. (2017). Early use of phonological codes in deaf readers: An ERP study. *Neuropsychologia*, 106, 261–279. <https://doi.org/10.1016/j.neuropsychologia.2017.10.006>
- Heathcote, A., Popiel, S. J., & Mewhort, D. J. (1991). Analysis of response time distributions: An example using the Stroop task. *Psychological Bulletin*, 109, 340–347. <https://doi.org/10.1037/0033-2909.109.2.340>
- Jacobs, A. M., Grainger, J., & Ferrand, L. (1995). The incremental priming technique: A method for determining within-condition priming effects. *Perception & Psychophysics*, 57, 1101–1110. <https://doi.org/10.3758/BF03208367>
- Ktori, M., Grainger, J., Dufau, S., & Holcomb, P. J. (2012). The “electrophysiological sandwich”: A method for amplifying ERP priming effects. *Psychophysiology*, 49, 1114–1124. <https://doi.org/10.1111/j.1469-8986.2012.01387.x>
- Ktori, M., Kingma, B., Hannagan, T., Holcomb, P. J., & Grainger, J. (2014). On the time-course of adjacent and non-adjacent transposed-letter priming. *Journal of Cognitive Psychology*, 26, 491–505. <https://doi.org/10.1080/20445911.2014.922092>
- Lenth, R., Singmann, H., Love, J., Buerkner, P., & Herve, M. (2020). *emmeans: Estimated marginal means* (R Package Version 1.4.8). <https://CRAN.R-project.org/package=emmeans>
- Lété, B., & Fayol, M. (2013). Substituted-letter and transposed-letter effects in a masked priming paradigm with French developing readers and dyslexics. *Journal of Experimental Child Psychology*, 114, 47–62. <https://doi.org/10.1016/j.jecp.2012.09.001>
- Lupker, S. J., & Davis, C. J. (2009). Sandwich priming: A method for overcoming the limitations of masked priming by reducing lexical competitor effects. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 35, 618–639. <https://doi.org/10.1037/a0015278>
- Lupker, S. J., Spinelli, G., & Davis, C. J. (2020a). Is zjudge a better prime for JUDGE than zudge is? A new evaluation of current orthographic coding models. *Journal of Experimental Psychology: Human Perception and Performance*, 46(11), 1252–1266. <https://doi.org/10.1037/xhp0000856>
- Lupker, S. J., Spinelli, G., & Davis, C. J. (2020b). Masked form priming as a function of letter position: An evaluation of current orthographic coding models. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 46, 2349–2366. <https://doi.org/10.1037/xlm0000799>
- Lupker, S. J., Zhang, Y. J., Perry, J. R., & Davis, C. J. (2015). Superset versus substitution-letter priming: An evaluation of open-bigram models. *Journal of Experimental Psychology: Human Perception and Performance*, 4, 138–151. <https://doi.org/10.1037/a0038392>
- Meade, G., Grainger, J., Midgley, K. J., Holcomb, P. J., & Emmorey, K. (2020). An ERP investigation of orthographic precision in deaf and hearing readers. *Neuropsychologia*,

- 146, 107542. <https://doi.org/10.1016/j.neuropsychologia.2020.107542>
- Monahan, P. J., Fiorentino, R., & Poeppel, D. (2008). Masked repetition priming using magnetoencephalography. *Brain and Language*, 106, 65–71. <https://doi.org/10.1016/j.bandl.2008.02.002>
- Ollman, R. T. (1977). Choice reaction time and the problem of distinguishing task effects from strategy effects. In S. Domic (Ed.), *Attention & performance VI* (pp. 99–113). Erlbaum.
- Perea, M., Abu Mallouh, R., & Carreiras, M. (2014). Are root letters compulsory for lexical access in Semitic languages? The case of masked form-priming in Arabic. *Cognition*, 132, 491–500. <https://doi.org/10.1016/j.cognition.2014.05.008>
- Perea, M., Marcet, A., Lozano, M., & Gomez, P. (2018). Is masked priming modulated by memory load? A test of the automaticity of masked identity priming in lexical decision. *Memory & Cognition*, 46, 1127–1135. <https://doi.org/10.3758/s13421-018-0825-5>
- Ratcliff, R. (1979). Group reaction time distributions and an analysis of distribution statistics. *Psychological Bulletin*, 86, 446–461. <https://doi.org/10.1037/0033-2909.86.3.446>
- R Core Team. (2020). *R: A language and environment for statistical computing*. R Foundation for Statistical Computing. <http://www.R-project.org/>
- Ridderinkhof, K. R., van den Wildenberg, W. P. M., Wijnen, J., & Burle, B. (2004). Response inhibition in conflict tasks is revealed in delta plots. In M. I. Posner (Ed.), *Cognitive neuroscience of attention* (pp. 369–377). The Guilford Press.
- Stinchcombe, E. J., Lupker, S. J., & Davis, C. J. (2012). Transposed-letter priming effects with masked subset primes: A re-examination of the “relative position priming constraint.” *Language and Cognitive Processes*, 27, 475–499. <https://doi.org/10.1080/01690965.2010.550928>
- Taikh, A., & Lupker, S. J. (2020). Do visible semantic primes preactivate lexical representations? *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 46(8), 1533–1569. <https://doi.org/10.1037/xlm0000825>
- Townsend, J. T., & Ashby, F. G. (1983). *Stochastic modeling of elementary psychological processes*. Cambridge University Press.
- Trifonova, I. V., & Adelman, J. S. (2018). The sandwich priming paradigm does not reduce lexical competitor effects. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 44, 1743–1764. <https://doi.org/10.1037/xlm0000542>
- Ziegler, J. C., Bertrand, D., Lété, B., & Grainger, J. (2014). Orthographic and phonological contributions to reading development: Tracking developmental trajectories using masked priming. *Developmental Psychology*, 50, 1026–1036. <https://doi.org/10.1037/a0035187>

