



Beyond visual similarity: Reversal-related masked priming in Korean Hangul

Sungbong Bae^{a,*}, Chang H. Lee^b, Manuel Perea^{c,d}

^a Department of Psychology, Yeungnam University, South Korea

^b Department of Psychology, Sogang University, South Korea

^c Departamento de Metodología and ERI-Lectura, Universitat de Valencia, Valencia, Spain

^d Centro de Investigación Nebrija en Cognición (CINC), Universitat Nebrija, Madrid, Spain

ARTICLE INFO

Keywords:

Reversal-related priming
Masked priming
Hangul
Orthographic processing
Orientation coding
Same–different matching

ABSTRACT

Reading requires precise encoding of letter identity to discriminate visually similar but orthographically distinct letter forms. We examined whether reversal-related alternatives (e.g., b and d) are more disruptive than other visually similar alternatives and whether any disadvantage differs between lateral and vertical reversal orientations. We conducted four masked priming experiments in Korean Hangul that examined lateral (e.g., ㅏ/ㅑ) and vertical (e.g., ㅓ/ㅕ) reversal-related primes in lexical decision, alphabetic decision, and same–different matching tasks, using stroke-deletion and stroke-addition controls and, in the final experiment, a visually dissimilar baseline. In the lexical decision task, reversal-related primes were more disruptive than both controls. The alphabetic decision task showed the same pattern relative to both visually similar controls, again with little evidence that these contrasts differed between orientations. In the same–different matching task, the disadvantage was larger for vertical than lateral pairs and, relative to the visually dissimilar baseline, remained for vertical but not lateral pairs. These findings extend the reversal-related disadvantage previously observed in Roman script to letter and word recognition in Hangul and show that reversal-related and visually similar single-stroke alternatives are not equivalent during Hangul word recognition. Thus, models of visual word recognition must distinguish reversal-related alternatives from visually similar alternatives formed by stroke deletion or addition.

Introduction

The human visual system can recognize objects across a series of transformations (Riesenhuber & Poggio, 1999). Its tolerance for left–right reflection (e.g., ㄹ and ㄿ) is often referred to as mirror generalization. We use the term in this sense and distinguish it from tolerance to picture-plane rotation; mirror-image invariance does not necessarily extend to rotations (see Bornstein et al., 1978; Fernandes & Leite, 2017; Gregory et al., 2011; Kolinsky et al., 2011; Logothetis & Pauls, 1995; Martinaud et al., 2016, for neuropsychological and developmental evidence). While adaptive for object perception, this ability conflicts with reading in the Roman alphabet, where letters with mirrored forms (e.g., b versus d, p versus q) represent different letters (Dehaene & Cohen, 2011). Thus, for readers of the Roman alphabet, literacy acquisition requires the cognitive system to become increasingly sensitive to orientation differences among letters, while preserving mirror tolerance

in object recognition (Dehaene, Nakamura, et al., 2010; see also Dehaene et al., 2015; Pegado et al., 2011).

Mirror errors are common in beginning readers (e.g., Fernandes et al., 2016) but typically decrease with literacy. This pattern is consistent with the *neuronal recycling* hypothesis (Dehaene et al., 2005). According to this hypothesis, reading repurposes visual circuits originally evolved for object processing and increases sensitivity to letter orientation. However, it remains unclear how the resulting reversal-related effects should be characterized in skilled readers and whether they generalize across writing systems. The key question is whether reversal-related pairs (e.g., d versus b) differ from visually similar pairs (e.g., c versus e). If they do, models of visual word recognition must distinguish between the two, even when both are visually similar. This constraint would apply to activation-based models (e.g., Grainger & Jacobs, 1996; McClelland & Rumelhart, 1981) and Bayesian evidence-accumulation models of masked priming (e.g., Kinoshita & Norris,

* Corresponding author.

E-mail addresses: sbbae@yu.ac.kr (S. Bae), changhwan1930@sogang.ac.kr (C.H. Lee), mperea@uv.es (M. Perea).

2012; Norris & Kinoshita, 2008). The same issue concerns how literacy enables readers to distinguish reversal-related letter forms, with implications for reading acquisition and dyslexia (e.g., Fernandes & Leite, 2017; Fernandes et al., 2024, 2025; Lachmann & Geyer, 2003).

To address these questions, we examined reversal-related priming in Hangul. Hangul contains systematic lateral (left–right; e.g., ㅏ/ㅑ) and vertical (top–bottom; e.g., ㅓ/ㅕ) reversal-related letter pairs and is a unicast script. This latter property avoids the case-dependent relations found in Roman letters: for example, lowercase *b* and *d* form a reversal pair, whereas uppercase *B* and *D* do not—a parallel case applies to *p* and *q* (*P/Q*). We tested whether reversal-related primes produce a cost relative to visually similar controls that are not reversal counterparts of the target and how masked priming from lateral and vertical reversal-related primes occurs in tasks that classify the target (lexical decision task, alphabetic decision task) or match it to a referent (same–different matching task). Specifically, using the masked priming technique (Forster & Davis, 1984), Experiments 1 and 2 employed lexical and alphabetic decision tasks, whereas Experiments 3 and 4 employed same–different matching tasks, in which participants compared the target with a preceding referent. We first review the relevant findings from the Roman alphabet.

Mirror and reversal-related processing in the Roman alphabet

A reversal-related cost has been reported in negative priming paradigms, where prior processing of reversible letters hinders mirror judgments. For instance, Borst et al. (2015) found that mirror-image animal targets were processed more slowly after reversible letter primes (e.g., *b/d*) than after nonreversible ones (e.g., *h/a*), suggesting that reversible letters can produce interference in subsequent mirror judgments. Likewise, Ahr et al. (2017) found a negative priming effect on subsequent object judgments after lateral but not vertical reversible-letter pairs. These findings, however, leave open the question of how reversal-related interference unfolds during the earliest moments of letter and word recognition.

To examine these early reversal-related effects, researchers have often used masked priming (Forster & Davis, 1984), which taps the initial moments of orthographic processing (see Grainger, 2008, for review). Perea et al. (2011) compared identity primes (*idea-IDEA*), mirror-reversed letter primes (*ibea-IDEA*), and visually similar controls (*ilea-IDEA*). Mirror-reversed primes selectively produced slower response times when they contained reversible letters (e.g., *b↔d*), but not when they contained nonreversible letters (e.g., *riɔe-RICE*), suggesting interference beyond visual similarity alone. Soares et al. (2019), using the same type of visually similar control (e.g., *lase-base*, *lose-dose*), further showed that the reversal-related cost within the *b/d* pair was asymmetric: reliable interference arose when a right-facing *b* prime replaced *d* in a *d*-word, but not in the reverse direction. Fernandes et al. (2022) found that mirror reflection and picture-plane rotation produced larger costs for reversible than nonreversible letters, with facilitation rather than interference for transformed nonreversible letters. Across these studies, then, a prime was selectively disruptive when its critical letter had an existing reversal counterpart.

Notably, visually similar letter-substitution primes that are not reversal counterparts of the target generally facilitate early word processing in masked priming experiments (e.g., *nevtral* vs. *neztral* → *NEUTRAL*; Marcet & Perea, 2017, 2018; Gutiérrez-Sigut et al., 2019; Lally & Rastle, 2023; see also Bae et al., 2025, for Hangul). This dissociation suggests that visual similarity per se cannot explain the reversal-related cost. Although spatially specified letter feature codes can distinguish reversal-related substitutions from other visually similar substitutions at the input level, this distinction does not by itself predict a selective disadvantage for reversal-related alternatives. Thus, the leading activation-based models that implement the Rumelhart and Siple (1974) letter-feature scheme can distinguish these substitutions at the input level, but they do not, by themselves, predict why relocating

a letter feature so as to specify a different letter identity produces greater interference than adding or deleting a letter feature (e.g., Coltheart et al., 2001; Davis, 2010; Grainger & Jacobs, 1996; McClelland & Rumelhart, 1981).

Most existing research on this issue has focused on laterally reversed letter pairs (e.g., *b–d*), leaving vertically reversed letter pairs (e.g., *b–p*) comparatively understudied. Given that mirror generalization specifically concerns tolerance to left–right reflection, vertical pairs provide a theoretically informative contrast, as both relations distinguish letter identities, but only lateral pairs instantiate the transformation implicated in mirror generalization. Previous developmental studies have shown that vertical confusion is less frequent and resolves earlier in reading development (Cairns & Steward, 1970; Davidson, 1935). For nonlinguistic objects, however, the relative prominence of the two reflection axes depends on visual experience and feature salience (see Cattaneo et al., 2010; Gregory & McCloskey, 2010). In skilled readers, Ahr et al. (2017) found that lateral reversible-letter pairs were harder to discriminate than vertical pairs. Notably, using whole-string masked primes in which every letter was transformed, Pittrich and Schroeder (2023) found facilitation relative to unrelated primes for left–right but not top–bottom transformations in the lexical decision task; in the same–different matching task, both transformations facilitated words and nonwords. As their manipulation transformed every letter, these findings provide relevant background for the lateral–vertical comparison but do not isolate the effect of a single reversal-related substitution.

A separate issue concerns the comparison baseline. In masked same–different matching experiments with single letters (e.g., *d–b*) and nonlexical strings (e.g., *dst–bst*), Kinoshita and Liong (2023) found that reversal-related primes facilitated responses relative to visually dissimilar controls but remained less effective than identity primes. They argued that the visually similar control letters used in previous experiments may themselves have facilitated the target—a question we return to in Experiment 4. Taken together, these findings leave two questions unresolved: whether, in masked word recognition, a single reversal-related substitution produces a cost relative to two visually similar alternatives, and how masked priming from lateral and vertical reversal-related primes is expressed in experimental tasks that classify the target (lexical decision, alphabetic decision) or match it to a referent (same–different matching task).

Hangul as a test case for reversal-related priming

The Korean writing system, Hangul, contains a systematic set of letter pairs that involve lateral and vertical reversals. Among the ten basic vowel letters, four form lateral (left–right) pairs (ㅏ/ㅑ, ㅓ/ㅕ), four form vertical (top–bottom) pairs (ㅓ/ㅗ, ㅕ/ㅛ), and two are symmetric (ㅗ, ㅣ). This inventory permits a within-script comparison of the two reversal orientations. In Korean orthography, Hangul letters are arranged within spatially organized syllable blocks rather than as a strictly linear sequence, and vowel letters do not occur as free-standing orthographic units. Within a syllable block, lateral vowels such as ㅏ and ㅑ occur to the right of the onset consonant (e.g., ㅏ + ㅑ = ㅑㅏ /sa/ [four]), whereas vertical vowels such as ㅓ and ㅗ occur below it (e.g., ㅏ + ㅓ = ㅓㅏ /so/ [cow]). For each reversal-related pair, however, one letter can be transformed into the other both by the relevant reflection (left–right or top–bottom) and by a 180° picture-plane rotation. The present materials therefore do not dissociate reflection from rotation, and the data speak most directly to reversal-related processing rather than to mirror-image invariance per se. We use *reversal-related* as the superordinate term for the Hangul contrasts and reserve *mirror* for cases where left–right reflection is specifically intended.

In addition to reversal-related letters, Hangul permits systematic contrasts based on stroke composition. Many letters derive from a base form by adding or removing a stroke (e.g., ㅓ–ㅗ–ㅛ; ㅣ–ㅑ–ㅕ; ㅓ–ㅗ–ㅛ). This structure allows direct comparison between reversal-related alternatives (e.g., ㅏ/ㅑ) and visually similar alternatives that are not reversal

counterparts of the target (e.g., ㅏ or ㅑ for ㅓ). Masked priming experiments in Korean Hangul have shown that primes missing a stroke tend to facilitate recognition of targets containing that stroke (e.g., ㅏ → ㅑ in ㅑ션-ㅑ션 [ku.fʌn/-/kʰu.fʌn/]), whereas the reverse does not hold (ㅑ → ㅏ in ㅑ독-ㅑ독 [kʰu.dok/-/ku.dok/]; Bae et al., 2025). This asymmetry parallels findings in Roman scripts, where primes with omitted diacritics facilitate recognition (e.g., facil → FÁCIL [easy]), but added-diacritic primes fail to do so (e.g., feliz → FELIZ [happy]; Perea et al., 2020; see Kinoshita et al., 2021, for an account of this asymmetry; see also Benyhe et al., 2023; Kinoshita et al., 2023; Solaja & Perea, 2026; Yu et al., 2023).

Orientation may also play a role. In the initial letter-pair discrimination task (a same-different judgment) of a negative-priming experiment, Winkler and Kim (2021) found that native Korean readers were slower to discriminate lateral reversal-related pairs (e.g., ㅏ + ㅑ) than vertical pairs (e.g., ㅏ + ㅓ), paralleling the lateral discrimination disadvantage that Ahr et al. (2017) reported with Roman letters. This lateral discrimination disadvantage raises the question of whether lateral and vertical primes also modulate the effects with the masked priming technique.

Thus, although reversal-related effects have primarily been studied with Roman letters, it remains unclear whether comparable patterns operate in Hangul. The two scripts provide different contexts for studying these effects: reversal relations in the Roman alphabet are case-dependent (e.g., d/b but not D/B), whereas Hangul is unicas and contains systematic lateral and vertical reversal-related letter pairs. Because these pairs distinguish both letter and word identities (e.g., 사 /sa/ [four] vs. 서 /sa/ [west]; 소 /so/ [cow] vs. 수 /su/ [number]), accurate discrimination between reversal-related letters is important for Korean letter and word recognition. In sum, the present experiments examined whether the reversal-related costs reported in the Roman alphabet extend to Hangul and what the results imply for models of visual word recognition.

Rationale of the experiments

To address these questions, we conducted four masked priming experiments in Hangul. The primary question was whether reversal-related primes would produce a cost relative to two visually similar controls that were not reversal counterparts of the target: Feature- (stroke deletion) and Feature+ (stroke addition). Experiment 1 tested this question in a lexical decision task using words presented in the standard Hangul format (i.e., words presented as syllable blocks). If reversal-related primes were more disruptive than both (visually similar) control conditions, this would indicate that reversal-related and visually similar alternatives are not functionally equivalent. The comparison between lateral and vertical pairs also assessed whether the reversal-related cost in lexical decision differed between the two reversal orientations.

Experiment 2 used an alphabetic decision task with isolated letters, in which the response classified the target as a letter or a pseudoletter, to examine whether the reversal-related disadvantage observed in word recognition would also occur in this letter-level decision context (see Grainger, 2024, for claims that this task parallels the lexical decision task at the letter level). Experiments 3 and 4 used a same-different matching task, in which the response concerned the correspondence between a preceding referent and the target. In the present design, the referent and target on “same” trials were the same glyph. Given that Hangul is unicas, physical-form identity and abstract letter identity coincided on these trials, so the present design does not allow us to determine which representational level primarily supports the match (see Carreiras et al., 2012; Kinoshita & Kaplan, 2008; Norris &

Kinoshita, 2008). Experiment 3 used the same visually similar Feature- and Feature+ controls as Experiment 2, whereas Experiment 4 replaced them with a visually dissimilar control to test whether any pattern obtained with the similar controls persisted relative to a visually dissimilar baseline. All four experiments also compared lateral and vertical pairs. Experiment 1 served as the primary test of Hangul word recognition; Experiments 2–4 examined reversal-related priming in single-letter classification and matching tasks.

Experiment 1 (lexical decision task)

To examine whether reversal-related primes affect lexical decisions during the early moments of visual word recognition, we crossed four prime conditions (Identity, Reversal-related, Feature-, and Feature+) with Reversal Orientation (lateral vs. vertical). The critical test was whether reversal-related primes (e.g., ㅏ ↔ ㅑ; ㅏ ↔ ㅓ) would produce slower lexical decisions than two visually similar controls: one differing by stroke deletion (Feature-) and one by stroke addition (Feature+).

Most prior experiments in the Roman alphabet compared three conditions: Identity (e.g., idea-IDEA), Reversal-related (e.g., ibea-IDEA with b as reversal-related to d), and a single visually similar non-reversal control prime (e.g., ilea-IDEA). The present Hangul design extended this comparison by including two visually similar controls—Feature- (stroke deletion, e.g., ㅏ [i] for ㅓ [a]) and Feature+ (stroke addition, e.g., ㅑ [ja] for ㅓ [a])—providing a finer-grained manipulation. This design allowed us to test whether any reversal-related disadvantage was limited to comparison with a stroke-deletion control or also occurred relative to a stroke-addition control.

Method

Data Availability. Materials, anonymized data, analysis scripts, and supplementary outputs for all four experiments are available at [https://osf.io/njp2u/overview?view_only=1ee2d77e673b4ec18934b82fa9426938].

Participants

Seventy-four undergraduate students (43 female, 31 male; $M = 22.40$ years, $SD = 1.32$, range: 19–27 years) from Yeungnam University participated for course credit. All were native Korean speakers with normal or corrected-to-normal vision. With 74 participants each completing 240 word trials (plus 240 nonword trials), the design yields $74 \times 60 = 4440$ observations per prime condition (given 4 prime conditions via Latin-square counterbalancing), which exceeds typical recommendations for masked priming experiments (e.g., Brysbaert & Stevens, 2018). All experiments were approved by the Institutional Review Board of Yeungnam University (Approval No. 2021-04-003-001) and conducted in accordance with the Declaration of Helsinki.

Materials and design

We selected 240 Korean disyllabic words of 4–5 Hangul letters ($M = 4.7$, $SD = 0.5$), all with an open (CV) first syllable (i.e., CV-CV or CV-CVC). The four critical letters formed one lateral reversal-related pair (ㅏ – ㅑ) and one vertical reversal-related pair (ㅏ – ㅓ); all four were vowels. Half of the target words contained a letter from the lateral pair as the first-syllable nucleus, and half contained a letter from the vertical pair in the same syllabic role. In the standard syllable-block format used for the word targets, letters in the lateral reversal pair appeared to the right of the onset consonant (e.g., 가, 7), whereas those in the vertical reversal pair appeared below it (e.g., ㅏ, 7). Control primes were created by replacing the critical letter with a letter from the set {ㅑ, ㅓ, ㅕ, ㅗ, ㅛ, ㅜ, ㅠ} that either lacked one stroke present in the

Table 1

Mean response times (in ms) and miss rates (%) for words used in Experiment 1.

Prime type	Lateral target words (e.g., 거짓)			Vertical target words (e.g., 구청)		
	Example	RT	Miss	Example	RT	Miss
Identity	거짓 – 거짓	541	4.3	구청 – 구청	536	3.2
Reversal-related	가짓 – 거짓	568	4.3	고청 – 구청	573	3.6
Feature–	기짓 – 거짓	551	4.1	그청 – 구청	548	3.2
Feature+	겨짓 – 거짓	557	4.5	규청 – 구청	554	3.7

target letter (Feature–) or contained one additional stroke (Feature+). We used a go/no-go masked priming lexical decision task with a 2×4 factorial design: Reversal Orientation (lateral vs. vertical) crossed with Prime Type (Identity, Reversal-related, Feature–, Feature+).

Mean word frequency per million (Kim, 2005) was 158 ($SD = 257$, range = 7–1331) for targets from the lateral pair and 166 ($SD = 218$, range = 8–1220) for targets from the vertical pair. The two target sets were closely matched in frequency, length, and orthographic neighborhood size (lateral: $M = 9.8$, $SD = 4.0$, range = 3–21; vertical: $M = 10.1$, $SD = 4.1$, range = 3–20). Mean visual similarity ratings for the critical prime–target letter pairs, obtained in an independent norming study (see Appendix A), were 4.69 for reversal-related primes, 4.08 for Feature– primes, and 4.89 for Feature+ primes. The corresponding values were 4.81, 4.48, and 4.93 for targets from the lateral pair, and 4.56, 3.68, and 4.85 for targets from the vertical pair.

For each target word, four prime types were generated: (a) Identity (e.g., lateral: 거짓–거짓; vertical: 구청–구청); (b) Reversal-related (replacing the critical letter with its reversal-related counterpart, always yielding a pseudoword; e.g., lateral: 가짓–거짓; vertical: 고청–구청); (c) Feature– (replacing the critical letter with a visually similar letter missing a stroke, always yielding a pseudoword; e.g., lateral: 기짓–거짓; vertical: 그청–구청); (d) Feature+ (replacing the critical letter with a visually similar letter formed by adding a stroke, always yielding a pseudoword; e.g., lateral: 겨짓–거짓; vertical: 규청–구청). This yielded 30 word targets per condition.

An additional 240 legal pseudowords were created by altering letters in existing words (e.g., 가섯, 두꼭). These pseudowords contained the same critical letters (ㄷ, ㅌ, ㅊ, ㅍ) in their initial syllables and were matched to words in length and neighborhood size (lateral: $M = 9.1$, $SD = 3.9$, range = 4–21; vertical: $M = 9.6$, $SD = 3.6$, range = 3–20). The prime–target manipulations for pseudoword trials paralleled those for words. Four counterbalanced lists were constructed via a Latin-square design: each target appeared once per participant but in a different prime condition. Participants were randomly assigned to one of the lists.

Procedure

Participants were tested individually in a quiet room using DMDX software (Forster & Forster, 2003) on a Windows computer connected to a monitor with a resolution of 1920×1080 pixels and a refresh rate of 60 Hz (16.7 ms per frame). They were seated approximately 60 cm from the screen. Participants were instructed to decide, as quickly and accurately as possible, whether each Hangul syllable string was a real word (by pressing the 'P' key) or a nonword (by withholding a response). They were not informed of the presence of a prime. Each trial consisted of three events: (a) a forward mask of four hash marks (####) for 500 ms; (b) a prime presented for 50 ms (3 frames); and (c) a target displayed until response or for a maximum of 1200 ms. To reduce visual overlap between primes and targets, primes appeared in a 20-point Gothic font, while targets were presented in the Batang font at a 20% larger size. Twenty practice trials with feedback preceded the experimental phase. No feedback was provided during the main task. The experimental phase comprised 480 trials (240 words, 240 nonwords), with short breaks every 120 trials, and lasted approximately 20 min.

Results and discussion

For the latency analyses, we excluded incorrect responses (i.e., failures to press “go” before the 1200-ms deadline; 3.8% of word trials) and RTs faster than 300 ms (0.1% of word trials). The mean correct response times (RTs) and miss rates are shown in Table 1.

RTs were analyzed with Bayesian linear mixed-effects models (brms; Bürkner, 2017) in R (R Core Team, 2021) after a $-1000/\text{RT}$ transformation under a Gaussian likelihood.¹ The model included fixed effects of Prime Type (Reversal-related = reference; Identity, Feature–, Feature+) and Reversal Orientation (lateral = -0.5 , vertical = $+0.5$). The random-effects structure was specified as $-1000/\text{RT} \sim \text{ReversalOrientation} \times \text{PrimeType} + (1 + \text{ReversalOrientation} \times \text{PrimeType} | \text{Subject}) + (1 + \text{PrimeType} | \text{Target})$, representing the maximal random-effects structure justified by the design, with varying intercepts and slopes for subjects and targets. The model was estimated with four chains and 5000 iterations (1000 warm-up). The four chains converged adequately (all $\hat{R} = 1.00$). The effects are reported as the posterior means (b), the standard deviation of the posterior distribution (SE), and the 95% credible interval (CrIs); effects are interpreted as credible when the 95% CrIs exclude zero. Accuracy was not analyzed further, given the low overall miss rate of 3.8% in the go/no-go task. The overall false-alarm rate on no-go (nonword) trials was 5.8%.

Reversal-related primes produced slower lexical decision times than identity primes ($b = -0.12$, $SE = 0.01$, 95% CrI [-0.13 , -0.10]; e.g., lateral: 거짓–거짓 < 가짓–거짓; vertical: 구청–구청 < 고청–구청). More importantly, target words were recognized more slowly following reversal-related primes than following either of the visually similar controls: Feature– ($b = -0.08$, $SE = 0.01$, 95% CrI [-0.09 , -0.06]; e.g., 기짓–거짓; 그청–구청) and Feature+ ($b = -0.06$, $SE = 0.01$, 95% CrI [-0.07 , -0.04]; e.g., 겨짓–거짓; 규청–구청) (see Table 1). The magnitude of this reversal-related cost did not credibly differ between the lateral and vertical pairs, as neither interaction was credible (Feature–: $b = -0.03$, $SE = 0.02$, 95% CrI [-0.06 , 0.00]; Feature+: $b = -0.02$, $SE = 0.02$, 95% CrI [-0.05 , 0.01]). The data did not provide credible evidence that the two visually similar controls differed from each other (Feature– vs. Feature+: $b = -0.02$, $SE = 0.01$, 95% CrI [-0.04 , 0.00]) or that this contrast differed between the lateral and vertical pairs ($b = -0.01$, $SE = 0.02$, 95% CrI [-0.04 , 0.02]).

In the present lexical decision experiment, reversal-related primes produced larger costs than both stroke deletion (Feature–; e.g., 기짓–거짓; 그청–구청) and stroke addition controls (Feature+; e.g., 겨짓–거짓; 규청–구청) in lexical decision. Thus, reversal-related alternatives showed a different pattern from visually similar non-reversal alternatives in Hangul word recognition.

Lexical decision alone, however, does not determine whether this effect should be attributed to whole-word processing, sublexical processing, or both. According to Grainger's (2024) hierarchical model of visual word recognition, letter processing forms an important foundation for lexical access, and the alphabetic decision task is particularly well suited to tap these processes. Specifically, this task provides a complementary decision context in which to examine whether the reversal-related cost also appears when the task places greater weight on single-letter categorization. Experiment 2 was designed to address this question.

¹ Given that the dependent measure in the analyses was $-1000/\text{RT}$, faster responses correspond to more negative values. Thus, with Reversal-related as the reference level, a negative coefficient for Identity or control-prime conditions indicates faster responses in that condition than in the Reversal-related condition. Following a reviewer's suggestion, we also conducted parallel analyses on raw RTs using an ex-Gaussian distribution for all four experiments. The critical pattern was unchanged; these analyses are reported on OSF.

Experiment 2 (alphabetic decision task)

Method

Participants

A separate sample of 104 Korean undergraduates (67 female, 37 male; $M = 22.73$ years, $SD = 1.04$) from Yeungnam University participated. Each completed 144 critical “go” trials, yielding a total of 14,976 data points (104×144). With eight experimental conditions, this provided 1872 observations per condition. All participants had normal or corrected-to-normal vision.

Materials and design

The task was a go/no-go masked priming alphabetic decision task. The experimental design and prime types (Identity, Reversal-related, Feature–, Feature+) matched those in Experiment 1, except that targets were individual Hangul letters rather than words. The targets were ㅏ and ㅑ for the lateral pair and ㅓ and ㅕ for the vertical pair, with equal numbers of trials in the two orientations.

For reversal-related primes, the target was preceded by its reversal counterpart (e.g., ㅏ for ㅑ, ㅓ for ㅕ). For Feature– controls, the target was preceded by a visually similar Hangul letter that lacks a stroke present in the target (e.g., ㅓ for ㅏ, ㅕ for ㅓ); for Feature+ controls, the target was preceded by a visually similar Hangul letter that contains an additional stroke (e.g., ㅑ for ㅏ, ㅕ for ㅓ). This yielded 16 unique prime–target combinations (4 targets $\times$ 4 prime types), each repeated 9 times, for 144 critical trials per participant (18 trials per condition). In addition, 48 filler trials using other Hangul letters (e.g., ㅑ, ㅓ, ㅕ, ㅗ, ㅛ, ㅜ, ㅠ) were included to broaden the stimulus set beyond the four critical targets and reduce potential reliance on a small repeated set.

Ten pseudoletters were constructed for no-go responses, each paired with a corresponding real Hangul letter and matched as closely as possible for size, stroke count, symmetry, junctions, and terminations. Given that no standardized pseudoletter set for Hangul comparable to BACS-2 (Vidal et al., 2017) exists, we developed a novel set following BACS-2 principles of attribute control; the full pseudoletter set is shown in Appendix B. The 192 no-go trials matched the 192 go trials (144 critical and 48 filler trials), yielding 384 trials in total.

Procedure

The trial sequence and apparatus were the same as in Experiment 1, except that participants decided whether the target was an existing Hangul letter or a pseudoletter, and the target remained on the screen until response or for a maximum of 1000 ms. They pressed the ‘P’ key for existing Hangul letters and withheld responses for pseudoletters.

Results and discussion

For the latency analyses on the letter targets, we excluded incorrect responses (i.e., failures to press “go” before the 1000-ms deadline; 0.2% of trials) and RTs faster than 250 ms (0.1% of trials; the lower bound was set at 250 ms rather than 300 ms to accommodate faster responses to single-letter targets). The mean RTs and the miss rates across conditions are presented in Table 2.

After these exclusions, RTs were analyzed using a Bayesian linear mixed-effects model specified as $-1000/RT \sim \text{ReversalOrientation} \times$

Table 2

Mean response times (in ms) and miss rates (%) for letters used in Experiment 2.

Prime type	Lateral target letter			Vertical target letter		
	Example	RT	Miss	Example	RT	Miss
Identity	ㅏ – ㅏ	402	0	ㅓ – ㅓ	392	0.2
Reversal-related	ㅑ – ㅏ	450	0.2	ㅕ – ㅓ	447	0.1
Feature–	ㅓ – ㅏ	436	0.1	ㅕ – ㅓ	429	0.2
Feature+	ㅑ – ㅏ	438	0.2	ㅕ – ㅓ	432	0.2

PrimeType + (1 + ReversalOrientation $\times$ PrimeType || Subject) + (1 | Target), with Prime Type and Reversal Orientation coded as in Experiment 1. The model included varying intercepts for targets and varying intercepts and slopes for subjects, with correlations among subject-level random effects suppressed to avoid convergence problems (Barr et al., 2013). All chains converged adequately ($R^2 = 1.00$). Given that miss rates were negligible, accuracy was not subjected to further analysis. The overall false-alarm rate on no-go pseudoletter trials was 0.6%.

As in Experiment 1, responses to target letters were slower following reversal-related primes than following identity primes ($b = -0.31$, $SE = 0.01$, 95% CrI $[-0.33, -0.28]$; e.g., lateral ㅏ – ㅏ < ㅑ – ㅏ; vertical ㅓ – ㅓ < ㅕ – ㅓ). Reversal-related primes also produced slower responses than the visually similar controls: Feature– ($b = -0.09$, $SE = 0.01$, 95% CrI $[-0.10, -0.07]$; e.g., ㅓ – ㅏ; ㅕ – ㅓ) and Feature+ ($b = -0.08$, $SE = 0.01$, 95% CrI $[-0.09, -0.07]$; e.g., ㅑ – ㅏ; ㅕ – ㅓ). The data did not provide credible evidence that either reversal-versus-control contrast differed between lateral and vertical pairs (Feature–: $b = -0.03$, $SE = 0.01$, 95% CrI $[-0.06, 0.00]$; Feature+: $b = -0.02$, $SE = 0.01$, 95% CrI $[-0.04, 0.01]$). The data also did not provide credible evidence that the two visually similar controls differed from each other (Feature– vs. Feature+: $b = -0.01$, $SE = 0.01$, 95% CrI $[-0.02, 0.01]$) or that this contrast differed between the lateral and vertical pairs ($b = -0.01$, $SE = 0.01$, 95% CrI $[-0.04, 0.01]$).

The results with the alphabetic decision task provided converging evidence for a reversal-related cost relative to both visually similar controls. Response times in the reversal-related conditions were consistently slower than in both the Feature– and Feature+ conditions (e.g., lateral: ㅏ – ㅏ = ㅑ – ㅏ < ㅓ – ㅏ; vertical: ㅓ – ㅓ = ㅕ – ㅓ < ㅑ – ㅓ), thereby paralleling the results of Experiment 1.

Experiment 3 tested the corresponding comparison in same–different matching, where participants judged whether a referent and target were the same.

Experiment 3 (same–different matching task)

Experiment 3 used same–different matching on isolated letters. As noted above, given that Hangul is unicas, physical-form identity and abstract letter identity coincided on “same” trials, so the experiment cannot determine by itself which representational level supported the match. Our question was whether reversal-related primes were more disruptive than Feature– and Feature+ controls under this discrimination environment, and whether this difference varied across reversal orientations.

Method

Participants

A total of 107 Korean undergraduates from Yeungnam University participated (62 female, 45 male; $M = 22.75$ years, $SD = 1.62$). Each participant provided 144 critical “same” trials for the analysis, yielding 15,408 observations in total (107×144), with 1926 observations per condition across the eight experimental conditions. All participants had normal or corrected-to-normal vision.

Materials and design

We employed a masked priming same–different matching task with a go/no-go procedure. The critical stimuli for “same” trials were identical to those in Experiment 2: lateral reversal-related letters (ㅏ, ㅑ) and vertical reversal-related letters (ㅓ, ㅕ). These letters were used as targets and their corresponding reversal-related letter primes. Control primes, identical to those in Experiment 2, were visually similar to the targets, differing by the absence or addition of a single stroke (e.g., ㅓ or ㅕ for ㅏ). The critical trials were evenly split between “same” and “different” response conditions, with 144 trials allocated to each. This yielded 18 “same” trials per condition. For critical “different” trials, the target was again one of the four critical letters {ㅏ, ㅑ, ㅓ, ㅕ}, whereas the referent

was a nonmatching letter. The 144 critical “different” trials comprised 48 reversal-related pairs, 48 feature-based pairs (24 with a Feature– referent and 24 with a Feature+ referent, defined relative to the target), and 48 cross-orientation pairs. Prime type (Identity, Reversal-related, Feature–, Feature+) was defined relative to the target, as on “same” trials. Each prime type occurred equally often on “same” and “different” trials and therefore did not predict the required response. Following the same rationale as in Experiment 2, an additional 96 filler trials (48 “same” and 48 “different”) were included, resulting in 384 trials in total.

Procedure

Participants completed the masked priming same–different matching task individually in a quiet room. They were instructed to view two sequentially presented letters—the referent and the target—and decide whether the two letters matched. For “same” trials, participants pressed the ‘P’ key; for “different” trials, no response was required (go/no-go procedure).

Each trial began with the presentation of the referent letter above a forward mask consisting of four hash marks (####) for 500 ms. The forward mask was then replaced by a prime stimulus, which was presented for 50 ms. Immediately following the prime, the target letter appeared and remained on the screen until the participant responded or for up to 1000 ms. Apart from these task-specific details, the procedure and apparatus were identical to those used in Experiments 1 and 2.

Results and discussion

For the latency analyses, we excluded incorrect responses on “same” trials (i.e., failures to press “go” before the 1000-ms deadline; 1.1% of trials) and RTs faster than 250 ms (0.4% of trials). Analyses were restricted to “same” trials, as “different” trials corresponded to no-go responses. The mean RTs and miss rates for each condition are summarized in Table 3. The data were analyzed using Bayesian linear mixed-effects modeling with the same specifications as in Experiment 2, and all chains converged adequately ($R^2 = 1.00$). Because miss rates were very low, no further analyses of accuracy were conducted. The overall false-alarm rate on no-go (“different”) trials was 1.8%.

As in Experiments 1 and 2, responses on “same” trials were slower following reversal-related primes than following identity primes ($b = -0.36$, $SE = 0.01$, 95% CrI $[-0.38, -0.33]$). They were also slower following reversal-related primes than following the visually similar Feature– and Feature+ controls: Feature– ($b = -0.25$, $SE = 0.01$, 95% CrI $[-0.28, -0.23]$) and Feature+ ($b = -0.11$, $SE = 0.01$, 95% CrI $[-0.13, -0.09]$).

Critically, in Experiment 3, both reversal-versus-control contrasts differed credibly between lateral and vertical pairs, with substantially larger costs for vertical pairs (reversal-related vs. Feature–: $b = -0.22$, $SE = 0.02$, 95% CrI $[-0.26, -0.19]$; reversal-related vs. Feature+: $b = -0.16$, $SE = 0.02$, 95% CrI $[-0.19, -0.12]$). In addition, the two visually similar controls differed from each other (Feature– vs. Feature+: $b = -0.15$, $SE = 0.01$, 95% CrI $[-0.17, -0.12]$), and this difference was modulated by orientation ($b = -0.07$, $SE = 0.02$, 95% CrI $[-0.10, -0.03]$), being larger for the vertical pair ($b = -0.18$, $SE = 0.02$, 95% CrI $[-0.21, -0.15]$) than for the lateral pair ($b = -0.11$, $SE = 0.02$, 95% CrI $[-0.14, -0.08]$).

In the present same–different matching experiment, both reversal-

versus-control contrasts were larger for vertical than for lateral pairs. Experiment 4 tested whether this lateral–vertical pattern would remain when the comparison baseline was visually dissimilar.

Experiment 4 (same–different matching with dissimilar controls)

Because the control primes in Experiment 3 were visually similar to the targets (Feature–: stroke deletion; Feature+: stroke addition), Experiment 3 alone could not determine whether the lateral–vertical difference arose from the reversal-related condition, the control conditions, or both.

To address this concern, in Experiment 4, we employed control primes that were markedly less visually similar to the targets than the Feature– and Feature+ controls, following the design rationale of Kinoshita and Liong (2023, Experiment 1). If reversal-related primes produce slower responses than dissimilar controls, this would constitute evidence of a reversal-related cost. This design tested whether the lateral–vertical pattern observed in Experiment 3 would also occur with a visually dissimilar baseline.

Method

Participants

Sixty-four Korean undergraduates from Yeungnam University participated (44 female, 20 male; $M = 22.36$ years, $SD = 1.28$). All were native Korean speakers with normal or corrected-to-normal vision. None had participated in the previous experiments.

Materials and design

The task and procedure were identical to those in Experiment 3 (masked priming same–different matching with a go/no-go procedure). The factorial design was 2 (reversal orientation: lateral vs. vertical) $\times$ 3 (prime type: Identity, Reversal-related, Dissimilar Control). The critical targets were the same four letters used in Experiments 2 and 3: ㅏ and ㅑ for the lateral pair and ㅓ and ㅕ for the vertical pair. The key difference from Experiment 3 was the use of a visually dissimilar control prime condition. Instead of the visually similar Feature– and Feature+ controls, each target was paired with a control prime that was markedly less similar in global configuration and was not a reversal-related counterpart of the target: ㅗ for lateral targets and ㅛ for vertical targets. Using ratings from the same norming study (see Appendix A), mean prime–target visual similarity was 4.69 for reversal-related primes (lateral: 4.81; vertical: 4.56) and 2.29 for dissimilar control primes (lateral: 2.56; vertical: 2.02). This yielded 12 unique prime–target combinations (4 targets $\times$ 3 prime types), each repeated 9 times, for a total of 108 “same” trials (18 trials per condition across 6 conditions). The 108 critical “different” trials comprised 36 reversal-related referent–target pairs and 72 cross-orientation pairs. Prime type (Identity, Reversal-related, Dissimilar) was defined relative to the target, as on “same” trials. Each prime type occurred equally often on “same” and “different” trials and therefore did not predict the required response. A total of 168 filler trials using other Hangul letters (ㅗ, ㅛ, ㅜ, ㅠ, ㅡ, ㅝ) were included to broaden the stimulus set beyond the four critical targets. These fillers were evenly divided into 84 “same” and 84 “different” trials, resulting in 384 trials in total.

Procedure

Participants completed a masked priming same–different matching task with a go/no-go procedure. All details of the procedure and apparatus were identical to Experiment 3.

Results and discussion

For the latency analyses, we excluded incorrect responses on “same” trials (misses; failures to press “go” before the 1000-ms deadline; 0.9% of trials) and RTs faster than 250 ms (0.2% of trials). Analyses were

Table 3
Mean response times (in ms) and miss rates (%) for letters used in Experiment 3.

Prime type	Lateral target letter			Vertical target letter		
	Example	RT	Miss	Example	RT	Miss
Identity	ㅏ – ㅏ	391	0.8	ㅓ – ㅓ	380	0.7
Reversal-related	ㅏ – ㅑ	419	1.4	ㅓ – ㅕ	470	2.5
Feature–	ㅏ – ㅗ	397	0.6	ㅓ – ㅛ	403	1.1
Feature+	ㅏ – ㅑ	414	0.7	ㅓ – ㅕ	431	0.9

Table 4
Mean response times (in ms) and miss rates (%) for letters used in Experiment 4.

Prime type	Lateral target letter			Vertical target letter		
	Example	RT	Miss	Example	RT	Miss
Identity	ㅏ - ㅏ	393	0.7	ㅓ - ㅓ	388	0.4
Reversal-related	ㅏ - ㅑ	423	0.9	ㅓ - ㅕ	471	2.7
Dissimilar	ㅏ - ㅓ	426	0.5	ㅏ - ㅓ	424	0.3

restricted to “same” trials, as “different” trials corresponded to no-go responses. The mean RTs and miss rates for each condition are summarized in Table 4. The data were analyzed using Bayesian linear mixed-effects modeling with the same specifications as in Experiments 2 and 3, and all chains converged adequately ($R^2 \leq 1.01$). Because miss rates were low, no further analyses of accuracy were conducted. The overall false-alarm rate on no-go (“different”) trials was 1.4%.

Overall, “same” responses were slower following reversal-related primes than following identity primes ($b = -0.33$, $SE = 0.02$, 95% CrI $[-0.36, -0.30]$) or dissimilar control primes ($b = -0.11$, $SE = 0.01$, 95% CrI $[-0.13, -0.08]$). Crucially, there was a credible interaction between reversal orientation and prime type, indicating that the reversal-related cost differed across orientations. This interaction was clear for both the identity–reversal contrast ($b = -0.27$, $SE = 0.02$, 95% CrI $[-0.32, -0.23]$) and the dissimilar-control–reversal contrast ($b = -0.26$, $SE = 0.02$, 95% CrI $[-0.30, -0.21]$), with both contrasts larger for vertical than for lateral pairs.

Planned simple-effect comparisons within each orientation clarified this interaction. For lateral targets, reversal-related primes did not credibly differ from dissimilar controls ($b = 0.02$, $SE = 0.02$, 95% CrI $[-0.01, 0.05]$), consistent with the negligible mean difference (Reversal-related: 423 ms; Dissimilar: 426 ms). In addition, reversal-related primes were slower than identity primes ($b = -0.20$, $SE = 0.02$, 95% CrI $[-0.23, -0.16]$; 30 ms). For vertical targets, reversal-related primes produced a sizable 47 ms cost relative to dissimilar controls ($b = -0.23$, $SE = 0.02$, 95% CrI $[-0.27, -0.20]$); they were also 83 ms slower than identity primes ($b = -0.47$, $SE = 0.02$, 95% CrI $[-0.51, -0.43]$). Although the rated similarity of the dissimilar controls was lower for vertical than for lateral targets (2.02 vs. 2.56, respectively), the mean RTs in the dissimilar-control condition were virtually identical across orientations (lateral: 426 ms; vertical: 424 ms). Thus, the vertical-specific reversal-related cost cannot be attributed simply to a slower dissimilar-control baseline for vertical items.

The results of Experiment 4 complement those of Experiment 3. Experiment 4 yielded credible Reversal Orientation $\times$ Prime Type interactions for both the identity–reversal and dissimilar-control–reversal contrasts. The latter extends the lateral–vertical pattern of Experiment 3 to a visually dissimilar baseline: the vertical reversal-related cost remained substantial, whereas the reversal-related and dissimilar conditions did not credibly differ for the lateral pair, although reversal-related primes were slower than identity primes.

General discussion

We designed four masked priming experiments, using a variety of tasks (lexical decision, alphabetic decision, same–different matching) in Korean Hangul to address two related questions: (1) whether a single reversal-related substitution is more disruptive than visually similar Feature– and Feature+ substitutions, and (2) how masked priming from lateral and vertical reversal-related primes is expressed across decision contexts. Fig. 1 summarizes the results.

Summary of key findings

Experiment 1 showed that, in lexical decision, reversal-related primes produced sizable costs relative to both stroke-deletion (Feature–) and stroke-addition (Feature+) controls. For example, for Hangul words such as 거짓 and 구청, reversal-related primes (e.g., 가짓 and 고청) produced slower responses than visually similar controls formed by deleting or adding a stroke.

In Experiment 2, alphabetic decisions were likewise slower following reversal-related primes than following either Feature– or Feature+ primes, with no credible evidence that these contrasts differed between lateral and vertical pairs. This pattern provides converging evidence that the disadvantage found in Experiment 1 also occurs in a single-letter classification context.

Experiments 3 and 4 examined this pattern using same–different matching with isolated letters. In Experiment 3, both reversal-versus-control contrasts were larger for vertical than for lateral pairs. Because the Feature– and Feature+ controls were visually similar to the targets, Experiment 3 alone could not determine whether this difference arose from the reversal-related condition, the control conditions, or both. Experiment 4 tested this issue with a visually dissimilar control. The vertical cost relative to the dissimilar control remained substantial (47 ms), whereas reversal-related and dissimilar primes did not credibly differ for the lateral pair (–3 ms); lateral reversal-related primes were nevertheless 30 ms slower than identity primes. In sum, the lateral–vertical difference in the reversal-related cost was only credible in the same–different matching task. Experiments 1 and 2 did not yield credible evidence of comparable orientation modulation in the classification tasks (lexical decision and alphabetic decision tasks).

Theoretical implications

The primary theoretical implication is that reversal-related alternatives (e.g., ㅏ–ㅑ) are not functionally equivalent to visually similar Feature– or Feature+ alternatives (e.g., ㅏ–ㅓ, ㅑ–ㅕ) during Hangul letter and word recognition. This dissociation cannot be captured by a mere similarity account. Indeed, reversal-related primes fell between Feature– and Feature+ primes in rated similarity but produced greater disruption than either control. The relevant distinction is that the present manipulations alter different properties of the letter code. While Feature– and Feature+ substitutions change the inventory of strokes, a reversal-related substitution preserves the basic stroke inventory but changes their spatial arrangement so as to specify a different letter unit. Thus, any successful model that addresses the mapping between visual codes and letter units must retain information about where visual features occur, rather than representing only their presence or overall similarity.

Competition among confusable letter representations offers one possible account of this pattern (e.g., Interactive Activation model; McClelland & Rumelhart, 1981; see also Grainger & Jacobs, 1996). A reversal-related prime may activate a competing letter unit, delaying target identification. In the critical Hangul pairs, the prime and target share the same basic strokes, but the distinctive stroke occurs in a different location and specifies a different letter identity. However, the substituted letter in each of the reversal-related, Feature–, and Feature+ primes was an existing Hangul letter, so activation of a different letter identity cannot by itself explain why reversal-related primes were more disruptive. Likewise, in models that use the letter-feature scheme introduced by Rumelhart and Siple (1974), oriented line segments are represented at specific spatial locations. This code can distinguish

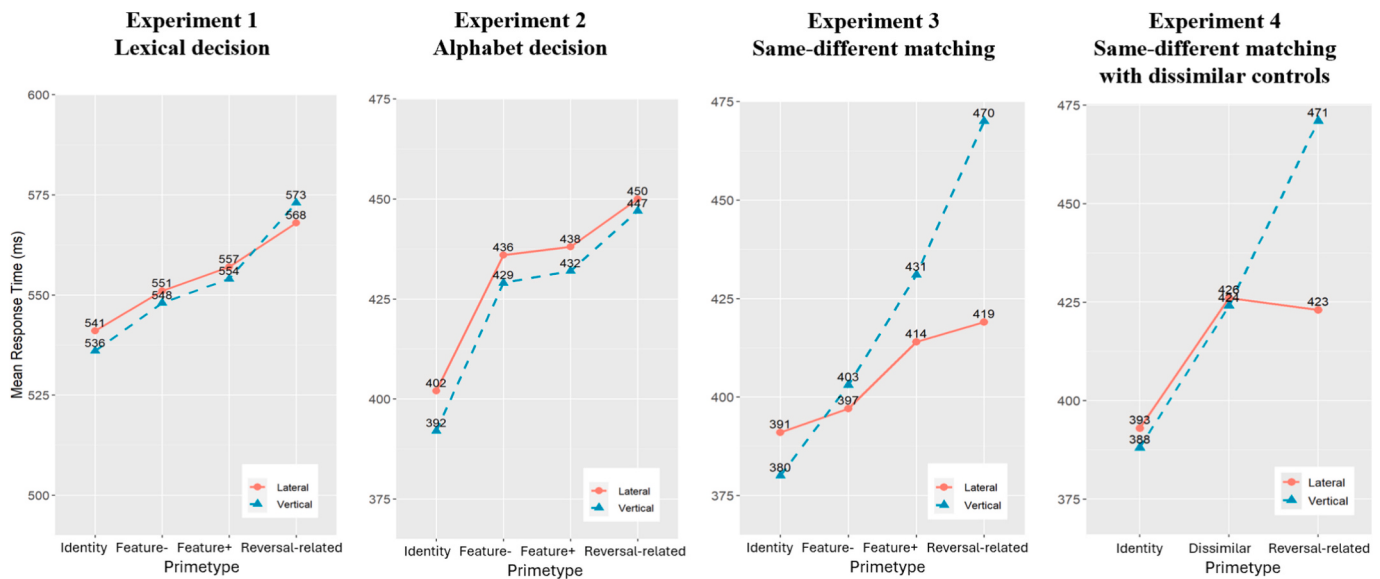


Fig. 1. Mean response times (ms) by prime type and reversal orientation (lateral vs. vertical). Panels from left to right show Experiments 1–4: lexical decision (Identity, Feature–, Feature+, Reversal-related), alphabetic decision (Identity, Feature–, Feature+, Reversal-related), same–different matching with similar controls (Identity, Feature–, Feature+, Reversal-related), and same–different matching with dissimilar controls (Identity, Dissimilar, Reversal-related). Points show condition means.

reversal-related substitutions from Feature– and Feature+ substitutions at the input level, but it does not by itself predict why reversal-related alternatives should produce greater interference than the single-stroke controls. The present data cannot determine whether this difference reflects stronger activation of a competing letter unit, lateral inhibition, or another mechanism. They do, however, constrain the mapping from spatially specified visual features onto letter identities. Any comprehensive model of visual word recognition must explain why rearranging the location of a stroke produces greater disruption than adding or deleting a stroke.

The reversal-related disadvantage observed in Experiment 1 extends previous findings from Roman-alphabet words (e.g., Perea et al., 2011; Soares et al., 2019) to Korean words written in Hangul. The change in writing system makes this a test of the cross-script generality of the representational distinction (cf. Simons et al., 2017). Hangul differs from the Roman alphabet in its letter forms, has no case alternation, and combines letters into syllable blocks. The Hangul result therefore shows that the distinction is not tied to the particular shapes of Roman reversal-related letters, case alternation, or the linear arrangement of letters in Roman-alphabet words. What the two writing systems share is the orthographic contrast: reversal-related forms can specify different letter identities. This pattern is consistent with literacy-based accounts of visual tuning (Dehaene, Pegado, et al., 2010; Pegado et al., 2011, 2014), according to which reading experience tunes letter identification to the contrasts needed to distinguish letter identities in a writing system. Accounts of early word recognition must therefore distinguish reversal-related substitutions from visually similar single-stroke substitutions, and this requirement is not limited to Roman-alphabet word forms. Its behavioral expression may nevertheless depend on script structure and the inventory of relevant contrasts (Winskel & Perea, 2018, 2022; see also Fernandes et al., 2021, for evidence from Tamil).

A separate, task-related implication concerns the lateral–vertical pattern across the four experiments. In the lexical and alphabetic decision tasks (Experiments 1 and 2), reversal-related primes produced greater disruption than Feature– and Feature+ primes, but there was no credible evidence that these reversal-versus-control contrasts differed by orientation. By contrast, in the same–different matching tasks (Experiments 3 and 4), the reversal-related cost was larger for the vertical pair than for the lateral pair; indeed, in Experiment 4, the vertical cost

relative to the dissimilar control was 47 ms, whereas there was no credible difference between the reversal-related and dissimilar-control conditions for the lateral pair. The Bayesian Reader account of masked priming provides a decision principle for interpreting this pattern (Kinoshita & Norris, 2012; Norris & Kinoshita, 2008). In this account, perceptual evidence accumulated from the masked prime and target is used to evaluate hypotheses defined by the task—whether the target is a word or a nonword, an existing Hangul letter or a pseudoletter, or a match to the referent. As each task evaluates a different hypothesis, the effect of the same prime–target relation may differ across tasks.² The Bayesian Reader account does not, by itself, predict this lateral–vertical pattern; here, its decision principle is used to connect prior discrimination findings to the decision required in each task.

The relevant findings come from the initial letter-pair stage of negative priming experiments using Roman letters in Ahr et al. (2017) and Hangul letters in Winskel and Kim (2021). In this stage, participants made a speeded same–different judgment, and lateral reversal-related pairs were harder to distinguish than vertical reversal-related pairs. Under the Bayesian Reader decision principle, this greater difficulty in distinguishing lateral pairs suggests that a lateral reversal-related prime in same–different matching may provide less evidence against the referent–target match than a vertical reversal-related prime. This could contribute to the smaller reversal-related cost for lateral than for vertical pairs. Lexical and alphabetic decision, however, do not require participants to compare the target with a referent: they classified it as a word or a nonword, or as an existing Hangul letter or a pseudoletter. The greater difficulty in distinguishing lateral pairs therefore need not produce a corresponding lateral–vertical difference in the reversal-related cost in these tasks. This interpretation is consistent with the larger vertical-than-lateral cost in Experiments 3 and 4 and with the absence of credible evidence that the reversal-versus-control contrasts differed by

² Task-dependent variation in masked priming effects is not specific to reversal-related primes. In Hangul, the asymmetric visual-similarity effect observed in lexical decision was absent in same–different word matching (Bae et al., 2025). Masked transposed-letter priming has likewise yielded different patterns in lexical decision and same–different matching tasks in Arabic (Boudelaa et al., 2019), Korean (Lee et al., 2021; Rastle et al., 2019), and Hebrew (Kinoshita et al., 2012; Velan & Frost, 2009).

orientation in Experiments 1 and 2. As the experiments did not directly compare task effects, this explanation is limited to the patterns observed within each task.

The preceding account suggests how the greater difficulty in distinguishing lateral than vertical reversal-related pairs may contribute to the task-related pattern, but it leaves open why lateral pairs are harder to distinguish. Ahr et al. (2017) proposed that this difference reflects visual experience. The same faces, animals, and objects are more often encountered in both left- and right-facing orientations than in both upright and upside-down orientations. They therefore suggested that mirror generalization may be stronger and harder to inhibit for lateral than for vertical reversals. Applied to the present findings, this proposal suggests that skilled readers may retain some tolerance for left–right reflection.³ This does not mean that they fail to distinguish the exemplars of the lateral pair as different letters; rather, lateral reversal-related pairs may remain more difficult to distinguish than vertical reversal-related pairs even after both letter contrasts have been learned. The present experiments did not test the source of this difference, so the visual-experience account remains a possible explanation.

Limitations and future directions

Several limitations should be noted. First, the critical reversal-related manipulations were restricted to vowels, in which Hangul offers the clearest lateral and vertical reversal-related contrasts. Given the empirical evidence that consonants and vowels can make different contributions to visual word recognition and lexical access (e.g., Berent & Perfetti, 1995; Carreiras et al., 2009; New et al., 2008), we acknowledge that the reversal-related effects might be modulated by consonant/vowel status. Second, for the letters used here, reflection and 180° picture-plane rotation map each member of a critical pair onto the same alternative. The present findings therefore cannot identify which transformation underlies the effects. Third, in Experiment 1, the critical letter was embedded within a Hangul syllable block, which may have modulated the magnitude of the reversal-related effects. Experiment 2 yielded a corresponding contrast with isolated letters, but this does not rule out modulation by syllabic grouping in Experiment 1. Finally, in Experiment 2, the four critical letters served as overt targets more often than the letters used as Feature– and Feature+ primes; thus, we cannot rule out the possibility that unequal overt-target exposure could have partly contributed to the reversal-versus-control contrasts.

Future research should extend the present design to include developing and less-skilled readers, as well as readers with atypical reading profiles. Automatic discrimination of mirrored and rotated letter relations appears to follow a protracted developmental trajectory rather than being fully established at the onset of literacy (see Fernandes et al., 2024), and mirror-related difficulties may remain detectable in dyslexic adult readers under masked priming (Fernandes et al., 2025). Time-sensitive measures such as event-related potentials could further establish when reversal-related costs emerge and how they evolve across decision contexts (Duñabeitia et al., 2011; Grainger & Holcomb, 2009). In addition, eye-movement experiments using the gaze-contingent boundary paradigm (Rayner, 1975) could examine whether reversal-related information is extracted parafoveally from upcoming words during sentence reading, extending the present findings to a more natural reading scenario. Future work should also use forms for which reflection and 180° picture-plane rotation yield different alternatives; 90°-rotation controls could additionally test whether the effects

generalize to other orientation changes (Bornstein et al., 1978; Logothetis & Pauls, 1995; Martinaud et al., 2016).

Conclusions

In a series of four experiments, we have shown that visual similarity alone cannot explain the reversal-related cost in masked priming. In both lexical and alphabetic decision, responses were slower following reversal-related primes than following visually similar controls. The greater disruption following reversal-related primes indicates that letter identification is sensitive to the spatial arrangement of visual features, rather than simply to their presence or overall similarity. The lexical decision findings extend the reversal-related disadvantage reported in the Roman alphabet to Hangul. A credible lateral–vertical difference emerged in the same–different matching task, whereas lexical and alphabetic decision tasks yielded no credible evidence of corresponding modulation. Overall, the findings show that reversal-related and visually similar single-stroke alternatives are not equivalent during letter and word recognition. Therefore, future computational models of visual letter and word recognition must distinguish reversal-related from visually similar letter substitutions.

CRediT authorship contribution statement

Sungbong Bae: Writing – original draft, Methodology, Formal analysis, Conceptualization. **Chang H. Lee:** Writing – review & editing, Methodology, Conceptualization. **Manuel Perea:** Writing – review & editing, Methodology, Formal analysis, Conceptualization.

Ethical statement

Informed consent and patient details

Written informed consent to take part in the study and to publish the article has been obtained from all participants or their legal representatives. The privacy rights of participants have been observed.

Studies in human

This study was performed in compliance with relevant laws, regulatory frameworks and guidelines where the research took place. This study was approved by the Institutional Review Board of Yeungnam University (Approval No. 2021–04–003–001)

Declaration of competing interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Acknowledgments

This work was supported by the Ministry of Education of the Republic of Korea and the National Research Foundation of Korea (NRF-2020S1A5B5A16083065), the Sogang University Research Grant (202415009.01 to Chang H. Lee), the Spanish Ministry of Science, Innovation, and Universities (PID2023-152078NB-I00), and the Valencian Government (Generalitat Valenciana; Grant CIAICO/2024/198).

³ The members of each critical pair are related both by reflection, left–right reflection for $\vdash-\dashv$ and up–down reflection for $\perp-\top$, and by 180° rotation. Within each pair, therefore, reflection cannot be dissociated from 180° rotation. However, given that the 180° rotational relation is shared by the lateral and vertical pairs, its presence alone cannot explain why the reversal-related cost was smaller for lateral than for vertical pairs. The observed lateral–vertical difference is consistent with greater residual tolerance for left–right reflection, but it does not uniquely establish left–right reflection as the source.

Appendix A. Visual Similarity Ratings for the Letter Pairs Used in the Experiments

The similarity ratings were obtained from a separate norming study, with a subset of the data previously reported in [Bae et al. \(2025\)](#). Fifty native Korean undergraduates, none of whom took part in the present experiments, completed one of two 800-trial lists that together covered all 1600 ordered pairs formed by the 40 Hangul letters. Letter order was counterbalanced across lists. Participants rated each pair on a 7-point scale (1 = *not at all similar*, 7 = *very similar*), following [Simpson et al. \(2013\)](#) and [Boudelaa et al. \(2020\)](#), and ratings were averaged across the two presentation orders. [Table A1](#) presents the means and standard deviations for the letter pairs used in Experiments 1–4.

Table A1
Mean Visual Similarity Ratings and Standard Deviations for the Letter Pairs Used in Experiments 1–4.

Condition	Orientation	Letter pair	<i>M</i>	<i>SD</i>
Reversal-related	Lateral	└─┘	4.81	1.48
Reversal-related	Vertical	┐┌┐	4.56	1.81
Feature−	Lateral	└─┘	4.61	1.48
Feature−	Lateral	└─┘	4.35	1.55
Feature−	Vertical	┐┌┐	3.54	1.57
Feature−	Vertical	┐┌┐	3.81	1.65
Feature+	Lateral	└─┘	5.06	1.27
Feature+	Lateral	└─┘	4.80	1.25
Feature+	Vertical	┐┌┐	4.73	1.28
Feature+	Vertical	┐┌┐	4.96	1.30
Dissimilar	Lateral	└─┘	2.52	1.53
Dissimilar	Lateral	└─┘	2.60	1.67
Dissimilar	Vertical	┐┌┐	1.89	1.09
Dissimilar	Vertical	┐┌┐	2.15	1.43

Appendix B. Hangul Letters and Pseudoletters

Hangul Letter	Pseudoletter
ㅏ	┌
ㅑ	┐
ㅓ	└
ㅕ	┘
ㅗ	┐
ㅛ	┐┐
ㅜ	┐┐
ㅠ	┐┐
ㅡ	┐┐
ㅣ	┐

Fig. B1. Hangul letter and pseudoletter stimuli. Each row presents a real Hangul letter (left) and its corresponding pseudoletter (right).

Data availability

The data, experimental materials, and analysis scripts are available at: https://osf.io/njp2u/?view_only=1ee2d77e673b4ec18934b82fa9426938

References

- Ahr, E., Houdé, O., & Borst, G. (2017). Predominance of lateral over vertical mirror errors in reading: A case for neuronal recycling and inhibition. *Brain and Cognition*, 116, 1–8. <https://doi.org/10.1016/j.bandc.2017.03.005>

- Bae, S., Lee, C. H., & Pae, H. K. (2025). Visual letter similarity effects in Korean word recognition: The role of distinctive strokes. *Quarterly Journal of Experimental Psychology*, 78(7), 1369–1378. <https://doi.org/10.1177/17470218241278600>
- 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(3), 255–278. <https://doi.org/10.1016/j.jml.2012.11.001>
- Benyhe, A., Labusch, M., & Perea, M. (2023). Just a mark: Diacritic function does not play a role in the early stages of visual word recognition. *Psychonomic Bulletin & Review*, 30(4), 1530–1538. <https://doi.org/10.3758/s13423-022-02244-4>
- Berent, I., & Perfetti, C. A. (1995). A rose is a REEZ: The two cycles model of phonology assembly in reading English. *Psychological Review*, 102(1), 146–184. <https://doi.org/10.1037/0033-295X.102.1.146>
- Bornstein, M. H., Gross, C. G., & Wolf, J. Z. (1978). Perceptual similarity of mirror images in infancy. *Cognition*, 6(2), 89–116. [https://doi.org/10.1016/0010-0277\(78\)90017-3](https://doi.org/10.1016/0010-0277(78)90017-3)

- Borst, G., Ahr, E., Roell, M., & Houdé, O. (2015). The cost of blocking the mirror generalization process in reading: Evidence for the role of inhibitory control in discriminating letters with lateral mirror-image counterparts. *Psychonomic Bulletin & Review*, 22(1), 228–234. <https://doi.org/10.3758/s13423-014-0663-9>
- Boudelaa, S., Norris, D., Mahfoudhi, A., & Kinoshita, S. (2019). Transposed letter priming effects and allographic variation in Arabic: Insights from lexical decision and the same-different task. *Journal of Experimental Psychology: Human Perception and Performance*, 45(6), 729–757. <https://doi.org/10.1037/xhp0000621>
- Boudelaa, S., Perea, M., & Carreiras, M. (2020). Matrices of the frequency and similarity of Arabic letters and allographs. *Behavior Research Methods*, 52(5), 1893–1905. <https://doi.org/10.3758/s13428-020-01353-z>
- Brysbaert, M., & Stevens, M. (2018). Power analysis and effect size in mixed effects models: A tutorial. *Journal of Cognition*, 1(1), 9. <https://doi.org/10.5334/joc.10>
- Bürkner, P. C. (2017). Brms: An R package for Bayesian multilevel models using Stan. *Journal of Statistical Software*, 80(1), 1–28. <https://doi.org/10.18637/jss.v080.i01>
- Cairns, N. U., & Steward, M. S. (1970). Young children's orientation of letters as a function of axis of symmetry and stimulus alignment. *Child Development*, 41(4), 993–1002. <https://doi.org/10.2307/1127327>
- Carreiras, M., Gillon-Downs, M., Vergara, M., & Perea, M. (2009). Are vowels and consonants processed differently? Event-related potential evidence with a delayed letter paradigm. *Journal of Cognitive Neuroscience*, 21(2), 275–288. <https://doi.org/10.1162/jocn.2008.21023>
- Carreiras, M., Perea, M., & Abu Mallouh, R. (2012). Priming of abstract letter representations may be universal: The case of Arabic. *Psychonomic Bulletin & Review*, 19, 685–690. <https://doi.org/10.3758/s13423-012-0260-8>
- Cattaneo, Z., Fantino, M., Silvano, J., Tinti, C., Pascual-Leone, A., & Vecchi, T. (2010). Symmetry perception in the blind. *Acta Psychologica*, 134(3), 398–402. <https://doi.org/10.1016/j.actpsy.2010.04.002>
- 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(1), 204–256. <https://doi.org/10.1037/0033-295X.108.1.204>
- Davidson, H. P. (1935). A study of the confusing letters b, d, p, and q. *The Pedagogical Seminary and Journal of Genetic Psychology*, 47(2), 458–468. <https://doi.org/10.1080/08856559.1935.10534056>
- Davis, C. J. (2010). The spatial coding model of visual word identification. *Psychological Review*, 117(3), 713–758. <https://doi.org/10.1037/a0019738>
- Dehaene, S., & Cohen, L. (2011). The unique role of the visual word form area in reading. *Trends in Cognitive Sciences*, 15(6), 254–262. <https://doi.org/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(4), 234–244. <https://doi.org/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(7), 335–341. <https://doi.org/10.1016/j.tics.2005.05.004>
- Dehaene, S., Nakamura, K., Jobert, A., Kuroki, C., Ogawa, S., & Cohen, L. (2010). Why do children make mirror errors in reading? Neural correlates of mirror invariance in the visual word form area. *NeuroImage*, 49(2), 1837–1848. <https://doi.org/10.1016/j.neuroimage.2009.09.024>
- Dehaene, S., Pegado, F., Braga, L. W., Ventura, P., Nunes Filho, G., Jobert, A., ... Cohen, L. (2010). How learning to read changes the cortical networks for vision and language. *Science*, 330(6009), 1359–1364. <https://doi.org/10.1126/science.1194140>
- Dunabeitia, J. A., Molinaro, N., & Carreiras, M. (2011). Through the looking-glass: Mirror reading. *NeuroImage*, 54(4), 3004–3009. <https://doi.org/10.1016/j.neuroimage.2010.10.079>
- Fernandes, T., Arunkumar, M., & Huettig, F. (2021). The role of the written script in shaping mirror-image discrimination: Evidence from illiterate, Tamil literate, and Tamil-Latin-alphabet bi-literate adults. *Cognition*, 206, Article 104493. <https://doi.org/10.1016/j.cognition.2020.104493>
- Fernandes, T., & Leite, I. (2017). Mirrors are hard to break: A critical review and behavioral evidence on mirror-image processing in developmental dyslexia. *Journal of Experimental Child Psychology*, 159, 66–82. <https://doi.org/10.1016/j.jecp.2017.02.003>
- Fernandes, T., Leite, I., & Kolinsky, R. (2016). Into the looking glass: Literacy acquisition and mirror invariance in preschool and first-grade children. *Child Development*, 87(6), 2008–2025. <https://doi.org/10.1111/cdev.12550>
- Fernandes, T., Pascual, M., & Araújo, S. (2025). Mirror invariance dies hard during letter processing by dyslexic college students. *Scientific Reports*, 15, 37395. <https://doi.org/10.1038/s41598-025-21092-5>
- Fernandes, T., Velasco, S., & Leite, I. (2024). Letters away from the looking glass: Developmental trajectory of mirrored and rotated letter processing within words. *Developmental Science*, 27(2), Article e13447. <https://doi.org/10.1111/desc.13447>
- Fernandes, T., Xavier, E., Domingues, M., & Araújo, S. (2022). Where is mirror invariance? Masked priming effects by mirrored and rotated transformations of reversible and nonreversible letters. *Journal of Memory and Language*, 127, Article 104375. <https://doi.org/10.1016/j.jml.2022.104375>
- Forster, K. I., & Davis, C. (1984). Repetition priming and frequency attenuation in lexical access. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 10(4), 680–698. <https://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(1), 116–124. <https://doi.org/10.3758/BF03195503>
- Grainger, J. (2008). Cracking the orthographic code: An introduction. *Language and Cognitive Processes*, 23(1), 1–35. <https://doi.org/10.1080/01690960701578013>
- Grainger, J. (2024). Letters, words, sentences, and reading. *Journal of Cognition*, 7(1), 66. <https://doi.org/10.5334/joc.396>
- 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(1), 128–156. <https://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(3), 518–565. <https://doi.org/10.1037/0033-295X.103.3.518>
- Gregory, E., Landau, B., & McCloskey, M. (2011). Representation of object orientation in children: Evidence from mirror-image confusions. *Visual Cognition*, 19(8), 1035–1062. <https://doi.org/10.1080/13506285.2011.610764>
- Gregory, E., & McCloskey, M. (2010). Mirror-image confusions: Implications for representation and processing of object orientation. *Cognition*, 116(1), 110–129. <https://doi.org/10.1016/j.cognition.2010.04.005>
- Gutiérrez-Sigut, E., Marcet, A., & Perea, M. (2019). Tracking the time course of letter visual-similarity effects during word recognition: A masked priming ERP investigation. *Cognitive, Affective, & Behavioral Neuroscience*, 19(4), 966–984. <https://doi.org/10.3758/s13415-019-00696-1>
- Kim, H. (2005). *Hyeondae gugeo sayong bindo josa 2 [a survey of the usage frequency of contemporary Korean 2]*. National Institute of the Korean Language.
- Kinoshita, S., Amos, A., & Norris, D. (2023). Diacritic priming in novice readers of diacritics. *Journal of Experimental Psychology: Human Perception and Performance*, 49(3), 370–383. <https://doi.org/10.1037/xhp0001084>
- Kinoshita, S., & Kaplan, L. (2008). Priming of abstract letter identities in the letter match task. *Quarterly Journal of Experimental Psychology*, 61(12), 1873–1885. <https://doi.org/10.1080/17470210701781114>
- Kinoshita, S., & Liong, G. (2023). Mirror letter priming is rightward-biased but not inhibitory: Little evidence for a mirror suppression mechanism in the recognition of mirror letters. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 49(10), 1523–1538. <https://doi.org/10.1037/xlm0001239>
- Kinoshita, S., & Norris, D. (2012). Task-dependent masked priming effects in visual word recognition. *Frontiers in Psychology*, 3, 178. <https://doi.org/10.3389/fpsyg.2012.00178>
- Kinoshita, S., Norris, D., & Siegelman, N. (2012). Transposed-letter priming effect in Hebrew in the same-different task. *Quarterly Journal of Experimental Psychology*, 65(7), 1296–1305. <https://doi.org/10.1080/17470218.2012.655749>
- Kinoshita, S., Yu, L., Verdonck, R. G., & Norris, D. (2021). Letter identity and visual similarity in the processing of diacritic letters. *Memory & Cognition*, 49(4), 815–825. <https://doi.org/10.3758/s13421-020-01125-2>
- Kolinsky, R., Verhaeghe, A., Fernandes, T., Mengarda, E. J., Grimm-Cabral, L., & Morais, J. (2011). Enantiomorphism through the looking glass: Literacy effects on mirror-image discrimination. *Journal of Experimental Psychology: General*, 140(2), 210–238. <https://doi.org/10.1037/a0022168>
- Lachmann, T., & Geyer, T. (2003). Letter reversals in dyslexia: Is the case really closed? *A critical review and conclusions. Psychology Science*, 45, 50–72.
- Lally, C., & Rastle, K. (2023). Orthographic and feature-level contributions to letter identification. *Quarterly Journal of Experimental Psychology*, 76(5), 1111–1119. <https://doi.org/10.1177/17470218221106155>
- Lee, C. H., Lally, C., & Rastle, K. (2021). Masked transposition priming effects are observed in Korean in the same-different task. *Quarterly Journal of Experimental Psychology*, 74(8), 1439–1450. <https://doi.org/10.1177/1747021821997336>
- 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>
- Marcet, A., & Perea, M. (2017). Is neutral NEUTRAL? Visual similarity effects in the early phases of written-word recognition. *Psychonomic Bulletin & Review*, 24(4), 1180–1185. <https://doi.org/10.3758/s13423-016-1180-9>
- Marcet, A., & Perea, M. (2018). Can I order a burger at ronalds.com? Visual similarity effects of multi-letter combinations at the early stages of word recognition. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 44(5), 699–706. <https://doi.org/10.1037/xlm0000477>
- Martinaud, O., Mirlink, N., Bioux, S., Blioux, E., Champmartin, C., Pouliquin, D., Cruyenninck, Y., Hannequin, D., & Gérardin, E. (2016). Mirrored and rotated stimuli are not the same: A neuropsychological and lesion mapping study. *Cortex*, 78, 100–114. <https://doi.org/10.1016/j.cortex.2016.03.002>
- 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(5), 375–407. <https://doi.org/10.1037/0033-295X.88.5.375>
- New, B., Araújo, V., & Nazzi, T. (2008). Differential processing of consonants and vowels in lexical access through reading. *Psychological Science*, 19(12), 1223–1227. <https://doi.org/10.1111/j.1467-9280.2008.02228.x>
- Norris, D., & Kinoshita, S. (2008). Perception as evidence accumulation and Bayesian inference: Insights from masked priming. *Journal of Experimental Psychology: General*, 137(3), 434–455. <https://doi.org/10.1037/a0012799>
- Pegado, F., Nakamura, K., Cohen, L., & Dehaene, S. (2011). Breaking the symmetry: Mirror discrimination for single letters but not for pictures in the visual word form area. *NeuroImage*, 55(2), 742–749. <https://doi.org/10.1016/j.neuroimage.2010.11.043>
- Pegado, F., Nakamura, K., & Hannagan, T. (2014). How does literacy break mirror invariance in the visual system? *Frontiers in Psychology*, 5, 703. <https://doi.org/10.3389/fpsyg.2014.00703>
- Perea, M., Fernández-López, M., & Marcet, A. (2020). What is the letter é? *Scientific Studies of Reading*, 24(5), 434–443. <https://doi.org/10.1080/10888438.2019.1689570>

- Perea, M., Moret-Tatay, C., & Panadero, V. (2011). Suppression of mirror generalization for reversible letters: Evidence from masked priming. *Journal of Memory and Language*, 65(3), 237–246. <https://doi.org/10.1016/j.jml.2011.04.005>
- Pittrich, K., & Schroeder, S. (2023). Priming effects in reading words with vertically and horizontally mirrored letters. *Quarterly Journal of Experimental Psychology*, 76(9), 2183–2196. <https://doi.org/10.1177/17470218221141076>
- R Core Team. (2021). R: A language and environment for statistical computing [Computer software]. R Foundation for Statistical Computing. <https://www.R-project.org/>.
- Rastle, K., Lally, C., & Lee, C. H. (2019). No flexibility in letter position coding in Korean. *Journal of Experimental Psychology: Human Perception and Performance*, 45(4), 458–473. <https://doi.org/10.1037/xhp0000617>
- Rayner, K. (1975). The perceptual span and peripheral cues in reading. *Cognitive Psychology*, 7(1), 65–81. [https://doi.org/10.1016/0010-0285\(75\)90005-5](https://doi.org/10.1016/0010-0285(75)90005-5)
- Riesenhuber, M., & Poggio, T. (1999). Hierarchical models of object recognition in cortex. *Nature Neuroscience*, 2(11), 1019–1025. <https://doi.org/10.1038/14819>
- Rumelhart, D. E., & Siple, P. (1974). Process of recognizing tachistoscopically presented words. *Psychological Review*, 81(2), 99–118. <https://doi.org/10.1037/h0036117>
- Simons, D. J., Shoda, Y., & Lindsay, D. S. (2017). Constraints on generality (COG): A proposed addition to all empirical papers. *Perspectives on Psychological Science*, 12(6), 1123–1128. <https://doi.org/10.1177/1745691617708630>
- Simpson, I. C., Mousikou, P., Montoya, J. M., & Defior, S. (2013). A letter visual-similarity matrix for Latin-based alphabets. *Behavior Research Methods*, 45(2), 431–439. <https://doi.org/10.3758/s13428-012-0271-4>
- Soares, A. P., Lages, A., Oliveira, H., & Hernández, J. (2019). The mirror reflects more for d than for b: Right asymmetry bias on the visual recognition of words containing reversal letters. *Journal of Experimental Child Psychology*, 182, 18–37. <https://doi.org/10.1016/j.jecp.2019.01.008>
- Solaja, O., & Perea, M. (2026). Are letter features weighted asymmetrically during the early moments of lexical access? Evidence from masked priming in Serbian Cyrillic. *Psychonomic Bulletin & Review*, 33, 231. <https://doi.org/10.3758/s13423-026-03006-2>
- Velan, H., & Frost, R. (2009). Letter transposition effects are not universal: The impact of transposing letters in Hebrew. *Journal of Memory and Language*, 61(3), 285–302. <https://doi.org/10.1016/j.jml.2009.05.003>
- Vidal, C., Content, A., & Chetail, F. (2017). BACS: The Brussels artificial character sets for studies in cognitive psychology and neuroscience. *Behavior Research Methods*, 49(6), 2093–2112. <https://doi.org/10.3758/s13428-016-0844-8>
- Winkel, H., & Kim, T.-H. (2021). The mirror generalization process in reading: Evidence from Korean Hangul. *Journal of Psycholinguistic Research*, 50(2), 447–458. <https://doi.org/10.1007/s10936-020-09736-1>
- Winkel, H., & Perea, M. (2018). Do the characteristics of the script influence responses to mirror letters? In N. Mani, R. K. Mishra, & F. Huettig (Eds.), *The interactive mind: Language, vision, and attention* (pp. 41–46). Macmillan.
- Winkel, H., & Perea, M. (2022). Mirror-image discrimination in monoliterate English and Thai readers: Reading with and without mirror letters. *Journal of Cultural Cognitive Science*, 6(2), 169–177. <https://doi.org/10.1007/s41809-021-00090-9>
- Yu, L., Zhang, Q., Ke, M., Han, Y., & Kinoshita, S. (2023). Some neighbors are more interfering: Asymmetric priming by stroke neighbors in Chinese character recognition. *Psychonomic Bulletin & Review*, 30(3), 1065–1073. <https://doi.org/10.3758/s13423-022-02207-9>