

Associative priming effects with visible, transposed-letter nonwords: JUGDE facilitates COURT

Manuel Perea · Dafna Palti · Pablo Gomez

Published online: 25 January 2012
© Psychonomic Society, Inc. 2012

Abstract Associative priming effects can be obtained with masked nonword primes or with masked pseudohomophone primes (e.g., *judpe*–COURT, *tode*–FROG), but not with visible primes. The usual explanation is that when the prime is visible, these stimuli no longer activate the semantic representations of their base words. Given the important role of transposed-letter stimuli (e.g., *jugde*) in visual word recognition, here we examined whether or not an associative priming effect could be obtained with visible transposed-letter nonword primes (e.g., *jugde*–COURT) in a series of lexical decision experiments. Results showed a sizable associative priming effect with visible transposed-letter nonword primes (i.e., *jugde*–COURT faster than *neevr*–COURT) in Experiments 1–3 that was close to that with word primes. In contrast, we failed to find a parallel effect with replacement-letter nonword primes (Experiment 2). These findings pose some constraints to models of visual word recognition.

Keywords Letter coding · Associative priming · Lexical decision

An important phenomenon for specifying the front-end of models of visual word recognition is the so-called *transposed-letter* effect: A nonword like *jugde* can be easily confusable with its base word, *judge* (Bruner & O’Dowd, 1958; O’Connor & Forster, 1981; Perea, Rosa, & Gómez, 2005; Rayner, White, Johnson, & Liversedge, 2006). The robustness of this effect poses some problems for slot-coding input schemes (i.e., the one used by the interactive activation model [McClelland & Rumelhart, 1981] and its successors) and has led to a cohort of more flexible input coding schemes (e.g., SERIOL model, Whitney, 2001; spatial coding model, Davis, 2010; overlap model, Gomez, Ratcliff, & Perea, 2008; open-bigram model, Grainger & van Heuven, 2003).

The focus of this report is to examine to what extent lexical/semantic activation from a transposed-letter nonword is sustained in time. To that end, we examine the impact of transposed-letter words in an associative priming paradigm. Prior studies have shown that when a transposed-letter nonword is presented briefly and masked, using a masked priming technique, there is access to associative information from the base word: Responses to the target COURT are faster when it is preceded by *jugde* than when it is preceded by the control nonword *neevr* (Perea & Lupker, 2003). Furthermore, Perea and Lupker (Experiment 1) found that the masked associative priming effect with word primes and with transposed-letter nonword primes was similar in magnitude (14.5 and 11 ms, respectively). Other masked priming studies have also shown similar priming effects from associatively mediated words (e.g., *tail*–STORY via *tale*, *tode*–FROG via *toad*, *judpe*–COURT via *judge*, or

M. Perea
Universitat de València,
Valencia, Spain

D. Palti
Bar-Ilan University,
Ramat-Gan, Israel

P. Gomez
DePaul University,
Chicago, USA

M. Perea (✉)
Departament de Metodologia, Facultat de Psicologia,
Av. Blasco Ibáñez, 21,
46010 València, Spain
e-mail: mperea@uv.es

toffee–CUP via *coffee*; see Bourassa & Besner, 1998; Drieghe & Brysbaert, 2002; Duñabeitia, Carreiras, & Perea, 2008; Lesch & Pollatsek, 1993; Lukatela & Turvey, 1994).

Importantly, in all the above-cited studies, the priming effect with an associatively mediated stimulus *vanished* when prime exposure duration allowed for the identification of the prime stimulus (see Bourassa & Besner, 1998; Drieghe & Brysbaert, 2002; Duñabeitia et al., 2008; Lesch & Pollatsek, 1993; Lukatela & Turvey, 1994). For instance, the one-letter-different nonword prime *judpe* facilitates the processing of COURT when it is presented very briefly (and masked), while it does not facilitate the processing of COURT when it is presented for 250 ms (Bourassa & Besner, 1998). The absence of a mediated associative priming effect with homophone/pseudohomophone/nonword primes has usually been interpreted in terms of an activation–verification account (Paap, Newsome, McDonald, & Schvaneveldt, 1982). At brief and masked presentations, the prime stimulus activates a number of potential candidates (i.e., *judpe* or *jugde* would activate *judge*, which in turn would preactivate COURT), while at long (conscious) exposures of the prime stimuli, the automatic activation that is obtained under masked priming conditions disappears. As Bourassa and Besner indicated, in the case of long prime exposures, “the nonword priming effect is eliminated because verification inhibits all activated lexical candidates (because they do not match the input)” (p. 66).

The main question we examine in the present article is whether the vanishing associatively mediated priming effect with visible primes also occurs with transposed-letter nonwords. The point here is that transposed-letter nonwords are perceptually very similar to the base words—even more so than one-letter-different neighbors (see Davis, 2010; Gomez et al., 2008). For instance, prior research has shown that the pattern of activation from the transposed-letter nonword prime decays more slowly than for a substitution-letter nonword prime (see, e.g., Perea et al., 2005, for evidence of a “pseudoword” frequency effect with transposed-letter nonwords, but not with substitution-letter nonwords). Furthermore, early in processing, ERP waves of transposed-letter nonwords are quite similar to the ERP waves of words (Carreiras, Vergara, & Perea, 2007).

In [Experiment 1](#), we conducted an unmasked lexical decision experiment using the Perea and Lupker (2003, Experiment 1) transposed-letter and word stimuli from their masked priming experiment. To examine whether the magnitude of the transposed-letter priming effect can be modulated by prime exposure duration, we chose two prime exposure durations: 200 and 400 ms. The rationale for this manipulation is that at a 200-ms prime exposure duration, there may be some remaining activation from the transposed-letter nonwords and, hence, associative priming may occur; at a 400-ms prime exposure duration, the “verification” stage is

presumably terminated, and associative information from the transposed-letter nonwords is less likely to have an effect in target processing. [Experiment 2](#) was analogous to [Experiment 1](#), except that we also included a related replacement-letter priming condition. Thus, this experiment provided a direct comparison between performance with word primes, transposed-letter primes, and replacement-letter primes (note that the materials in this experiment are the same as those in Perea and Lupker’s masked priming experiment). Finally, Experiment 3 was designed to provide a replication of [Experiment 1](#).

Experiment 1

Method

Participants A total of 40 students (20 at the 200-ms SOA and 20 at the 400-ms SOA) from the Massachusetts Institute of Technology took part in the experiment in exchange for a small monetary compensation. All of them had either normal vision or vision that was corrected to normal and were native speakers of English.

Materials The stimuli (120 word pairs and 120 nonword pairs) were taken from Perea and Lupker (2003, Experiment 1). For each word target, the prime was (1) a word associated to the target (associatively related word condition; e.g., *judge*–COURT), (2) a transposed-letter nonword created by transposing two internal letters from the associated prime (associatively related transposed-letter condition; e.g., *jugde*–COURT), (3) an unrelated word (unrelated word condition; e.g., *never*–COURT), or (4) an unrelated transposed-letter nonword (unrelated transposed-letter condition; e.g., *neevr*–COURT). Word primes and transposed-letter primes were counterbalanced throughout the related/unrelated conditions, so that each target word was primed by each of the four types of primes across the experiment. Four lists of materials were created. Different groups of participants were used for each list. For the purposes of the lexical decision task addition, we selected a set of 120 standard nonwords from Perea and Lupker’s [Experiment 1](#). These nonwords were preceded by 60 unrelated word primes (e.g., *thumb*–BRAMP) and 60 unrelated nonword primes (e.g., *shile*–KREMP).

Procedure The experiment was run individually in a quiet room. Presentation of the stimuli and recording of response times (RTs) were controlled by PC-compatible computers using DMDX (Forster & Forster, 2003). On each trial, a fixation point (+) was presented for 500 ms in the center of the screen. Next, the lowercase prime was presented for 200 or 400 ms, depending on the prime exposure duration

condition. The prime was immediately followed by the presentation of the target stimulus in uppercase. RTs were measured from target onset to the participant's response. All the strings were presented centered, in Courier New 12-point font. Participants were instructed to press the “M” button if the string formed an existing English word and the “Z” button if the string was a nonword. Each participant received a different order of trials. The whole experimental session lasted for about 14 min.

Results and discussion

Incorrect responses (4.9% of the data for word targets) and RTs less than 250 ms or greater than 1,500 ms (less than 1% of the data for word targets) were excluded from the latency analysis. The mean RTs and error percentages from the subject analyses are presented in Table 1. By-subject and by-item analyses of variance (ANOVAs) based on the participants' response latencies and percentages of errors were conducted on the basis of a 2 (associative relatedness: related or unrelated) $\times$ 2 (type of prime: word prime, transposed-letter nonword prime) $\times$ 2 (prime exposure duration: 200, 400 ms) $\times$ 4 (list: list 1, list 2, list 3, list 4) design. List was included as a factor in the design to extract the variance due to the error associated with the lists (see Pollatsek & Well, 1995).

The ANOVA on the latency data showed that responses to word targets were slower when they were preceded by a related prime than when they were preceded by an unrelated prime, $F_1(1, 32) = 18.76$, $p < .001$, $F_2(1, 116) = 25.58$, $p < .001$. This associatively related priming effect was similar in magnitude for word primes and transposed-letter nonword primes (17.5 vs. 18.5 ms, respectively), as can be

deduced by the lack of interaction between relatedness and type of prime (both $F_s < 1$). The associative priming effect was numerically smaller at the 400-ms prime exposure duration and at the 200-ms prime exposure duration, but the interaction between relatedness and prime exposure duration was not significant (both $p_s > .25$). The other effects were not significant (all $p_s > .20$), except for a main effect of prime exposure duration in the item analysis, $F_2(1, 116) = 139.78$, $p < .001$, $F_1(1, 32) = 2.38$, $p = .13$.

The ANOVA on the error data did not reveal any significant effects (all $p_s > .12$).

The results of the present experiment are clear: Associative priming effects can be observed with transposed-letter nonword primes (e.g., RTs to *jugde*–*COURT* were shorter than the RTs for *neevr*–*COURT*) even when the primes are clearly visible, thus extending the findings of Perea and Lupker (2003) with masked transposed-letter nonword primes. Indeed, the size of the associative priming effect for transposed-letter nonword primes was remarkably similar to that with word primes (18.5 and 17.5 ms, respectively); the parallel effects with masked primes in the Perea and Lupker experiment were 11 versus 14 ms, respectively. Finally, the semantic activation from the transposed-letter nonwords seems to be sustained in time; the magnitude of the associative priming effect for transposed-letter nonwords was similar when the prime exposure durations were 200 and 400 ms. Numerically, the associative priming effect was higher at the 200-ms prime exposure duration; however, leaving aside that the critical interaction did not approach significance, RTs were around 50 ms shorter for the group with the 400-ms prime exposure duration than for the group with the 200-ms prime exposure duration (note that shorter RTs tend to lead to smaller effects).¹

The present experiment shows that, unlike in previous experiments with word/nonword/pseudohomophone primes in which an associatively mediated priming effect disappeared when the prime was visible, the effect does *not* appear for transposed-letter nonword primes: Transposed-letter nonword primes produce associative priming effects not only when the primes are presented masked (Perea & Lupker, 2003), but also when the primes are clearly visible. Thus, the results support the view that transposed-letter neighbors do not behave as one-letter-different neighbors (see Duñabeitia, Perea, & Carreiras, 2009, for further evidence) and that transposed-letter nonwords have a very high degree of orthographic similarity with their base words—stronger than other types of nonwords (see Gomez et al., 2008; Perea & Fraga, 2006).

Table 1 Mean lexical decision times (in milliseconds) and percentages of errors (in parentheses) for word targets in Experiment 1

	Type of Prime		
	Related	Unrelated	Priming
200-ms prime exposure duration			
Words	670 (3.8)	691 (5.7)	21 (1.9)
TL nonwords	669 (5.0)	695 (5.7)	26 (0.7)
400-ms prime exposure duration			
Words	616 (4.7)	630 (3.9)	14 (-0.8)
TL nonwords	627 (4.8)	640 (6.0)	13 (1.2)
Replication 400-ms prime exposure duration			
Low-proportion related pairs			
Words	627 (2.6)	639 (5.4)	12 (2.8)
TL nonwords	640 (3.8)	651 (5.4)	11 (1.6)

¹ It may be important to note that the 13.5-ms associative priming effect at the 400-ms prime exposure duration was significant, $F_1(1, 16) = 5.90$, $p < .03$, $F_2(1, 116) = 10.42$, $p < .003$.

Contrary to our expectations, the magnitude of associative priming effects with transposed-letter nonword primes did not differ at the two prime exposure durations (200 and 400 ms). Indeed, the size of the effect was somewhat smaller at the longer prime exposure duration (note that same trend occurred for associatively related word primes). One could argue that perhaps the obtained associative priming effect, particularly at the 400-ms prime exposure duration, was affected by the participants noticing the prime–target relationships (e.g., “if prime and target are related, say ‘yes’”). To examine this issue, we conducted a replication of the 400-ms prime exposure duration subexperiment using a low proportion of related pairs (i.e., adding 240 filler, unrelated pairs).² Twenty MIT undergraduates participated in the experiment. Results showed an 11.5-ms associative priming effect, $F_1(1, 16) = 5.94, p < .03, F_2(1, 116) = 2.92, p = .090$, which was of similar magnitude for word primes and for transposed-letter nonword primes (12 vs. 11 ms, respectively; interaction effect: both $F_s < 1$; see Table 1). The ANOVA for the latency data also revealed a small cost of processing the transposed-letter stimuli, as deduced from the effect of type of prime, $F_1(1, 16) = 3.96, p = .065, F_2(1, 116) = 4.88, p < .03$. Thus, the pattern of data obtained in the 400-ms prime exposure duration in Experiment 1 was not an empirical anomaly. Instead, the small priming effects obtained at this long stimulus onset asynchrony (SOA) suggest that, on a number of trials, participants did not process the prime at a deep, semantic level.

The present experiment has one limitation, though. It did not test the associative priming effect for letter substitution primes (e.g., *judpe*–COURT vs. *nemer*–COURT). Experiment 2 was designed to directly examine associative priming effects for word primes (e.g., *judge*–COURT vs. *never*–COURT), transposed-letter nonword primes (*jugde*–COURT vs. *neevr*–COURT), and replacement-letter nonword primes (*judpe*–COURT vs. *nemer*–COURT) when the primes are visible. Indeed, Perea and Lupker (2003, Experiment 1) included these three conditions in their masked priming experiment. If the data from Experiment 2 reveal a significant associative priming effect for word primes and for transposed-letter nonword primes, but not for replacement-letter nonword primes (as actually happened in the Perea & Lupker masked priming experiment), this would add further evidence for the special role of transposed-letter stimuli in visual word recognition. In Experiment 2, prime exposure duration was set to 200 ms.

² From the 240 filler trials, there were 120 word trials (60 nonword–word trials and 60 word–word trials) and 120 nonword trials (60 nonword–nonword trials and 60 word–nonword trials). Thus, the proportion of related trials (among word targets) was .25; it was .50 in Experiment 1.

Experiment 2

Method

Participants A total of 54 students from DePaul University took part in the experiment in exchange for course credit. All of them had either normal vision or vision that was corrected to normal and were native speakers of English.

Materials The stimuli were the same as those in Experiment 1, except that we added two priming conditions for the word trials, which were also included in Perea and Lupker’s (2003) Experiment 1: (1) an associatively related replacement-letter condition, in which a replacement-letter nonword prime was created by replacing an internal letter from the associated prime (e.g., *judpe*–COURT), and (2) an unrelated replacement-letter condition, in which the prime was an unrelated replacement-letter nonword (e.g., *nemer*–COURT). Word primes, transposed-letter primes, and replacement-letter primes were counterbalanced throughout the related/unrelated conditions (i.e., each target word was primed by each of the six types of primes across the experiment). Six lists of materials were created. Different groups of participants were used for each list. The set of nonword targets was the same as that in Experiment 1. However, to keep the proportion of word/nonword primes, these nonwords were preceded by 40 unrelated word primes and 80 unrelated nonword primes; these were the same prime–target pairs as those in the Perea and Lupker study.

Procedure The procedure was the same as that in the 200-ms SOA condition in Experiment 1.

Results and discussion

Incorrect responses (4.1% of the data for word targets) and RTs less than 250 ms or greater than 1,500 ms (less than 0.86% of the data for word targets) were excluded from the latency analysis. The mean RTs and error percentages from the subject analyses are presented in Table 2. Subject and item ANOVAs based on the participants’ response latencies and percentages of errors were conducted on the basis of a 2 (associative relatedness: related or unrelated) $\times$ 3 (type of prime: word prime, transposed-letter nonword prime, replacement-letter nonword prime) $\times$ 6 (list: list 1, list 2, list 3, list 4, list 5, list 6) design.

The ANOVA for the latency data showed that responses to word targets were faster when they were preceded by a related prime than when they were preceded by an unrelated prime, $F_1(1, 48) = 21.80, p < .001, F_2(1, 114) = 19.44, p < .001$. The main effect of type of prime was not significant (both $p_s > .25$). More important, the magnitude of the

Table 2 Mean lexical decision times (in milliseconds) and percentages of errors (in parentheses) for word targets in [Experiment 2](#)

	Type of Prime		
	Related	Unrelated	Priming
Words	597 (3.1)	622 (4.2)	25 (1.1)
TL nonwords	606 (3.6)	617 (3.7)	11 (0.1)
RL nonwords	611 (4.2)	617 (5.9)	6 (1.7)

associatively related priming effect was modulated by the type of prime, as deduced by the significant interaction between associative relatedness and type of prime, $F_1(2, 96) = 4.21$, $p < .02$, $F_2(2, 228) = 3.06$, $p < .05$. The associative priming effect was rather large (25 ms) for the target words preceded by a word prime, $F_1(1, 48) = 18.21$, $p < .001$, $F_2(1, 114) = 22.81$, $p < .001$; it was smaller (11 ms) but statistically significant (in the analysis by subjects) for the target words preceded by a transposed-letter nonword prime, $F_1(1, 48) = 5.23$, $p < .03$, $F_2(1, 114) = 2.82$, $p = .096$, and it was even smaller (6 ms), and nonsignificant, for the target words preceded by a replacement-letter nonword prime, $F_1(1, 48) = 2.31$, $p = .135$, $F_2(1, 114) = 1.35$, $p > .20$.

The ANOVA for the error data revealed a significant effect of relatedness, $F_1(1, 48) = 6.82$, $p < .02$, $F_2(1, 114) = 4.18$, $p < .05$. The effect of type of prime was not significant (both $ps > .13$). Finally, the interaction between the two factors approached significance in the analysis by subjects, $F_1(2, 84) = 2.65$, $p = .077$, $F_2 < 1$, which reflected that the associative priming effect was numerically larger for replacement-letter primes than for the other two types of primes.

The main finding of the present experiment is that the magnitude of the associative priming effect was modulated by the type of prime. First, as in [Experiment 1](#), we found a significant associative priming effect for word primes and for transposed-letter nonword primes, although the size of the effect for transposed-letter nonwords was somewhat smaller than that in [Experiment 1](#) with the same prime duration (11 vs. 26 ms). Second, as in previous research, we failed to find a significant effect of associative priming effect for replacement-letter pseudoword primes. We should note that we found some hints of an effect (6 ms in the RT data and 1.7% in the error data); this suggests that, occasionally, participants could activate the meaning of the base word corresponding to the replacement-letter nonword prime. (We discuss this issue in the [General Discussion](#) section.) Thus, the present experiment provides further empirical evidence that transposed-letter pseudowords are perceptually closer to their base words than are replacement-letter pseudowords.

In the present experiments, the set of prime/target stimuli for the nonword trials was created merely for the purposes of the lexical decision task; they were taken from Perea and Lupker's (2003) [Experiment 1](#). That is, there was no degree of relationship for nonword pairs. One could argue that, in our experiments, the relationship between the prime and the nonword on nonword trials should have been the same as the one between the prime and the word on word trials. That is, one would need to have not only *judge*–*COURT* and *never*–*COURT*, but also *judge*–*BOURT* and *never*–*BOURT*, where *BOURT* would be a nonword created from the word *COURT*. In this way, the primes for the nonword trials would be similar to the primes for the word trials (i.e., words with highly associated prime stimuli). The rationale is that perhaps the participants could pick up the fact that one type of words (those with high associates) were always followed by a "yes" response, which in principle could explain (part of) the associative priming effect. However, this explanation cannot account for why there is an associative priming effect with the transposed letter nonwords—unless one assumes that transposed-letter nonwords have access to the semantic information of their base words. Furthermore, Perea and Rosa (2002) showed that at prime exposure durations (166 ms) very close to the one employed here (200 ms), the size of associative priming effects was not affected by the proportion of related pairs (i.e., an ideal "strategy" manipulation).

One apparent discrepancy between Experiments 1 and 2 is that the size of the associative priming effect for transposed-letter nonword primes was less robust in [Experiment 2](#) than in [Experiment 1](#) with the same prime exposure duration (11 vs. 26 ms), while it was remarkably similar for word primes (25 vs. 21 ms, respectively). One potential explanation for this difference for transposed-letter nonwords is that the inclusion of replacement-letter nonword primes may have led participants to a finer-grained processing of the prime stimuli (note that two thirds of the primes were nonwords, while in [Experiment 1](#) only half of the primes were nonwords). Another possibility is that the nature of associative/semantic priming effects varies across individuals. For instance, Yap, Tse, and Balota (2009) reported that the magnitude of semantic priming (and its interaction with frequency and nonword type) varied across participants from different universities, probably because of vocabulary knowledge and years of education in the participant pools. In our case, [Experiment 1](#) was conducted at MIT, and [Experiment 2](#) was conducted at DePaul University. We believed that it was important to reexamine the central finding of the present study—namely, the presence of associative priming with visible, transposed-letter nonword primes—with DePaul students, using the same conditions as in [Experiment 1](#) with a 200-ms prime exposure duration. This was the goal of Experiment 3.

Experiment 3 (replication of Experiment 1 with a 200-ms SOA)

Method

Participants Twenty-four students from DePaul University participated in the experiment in exchange for course credit. All of them had either normal vision or vision that was corrected to normal and were native speakers of English.

Materials and procedure The materials and procedure were the same as those in Experiment 1.

Results and discussion

Incorrect responses (5.1% of the data for word targets) and RTs less than 250 ms or greater than 1,500 ms (less than 0.97% of the data for word targets) were excluded from the latency analysis. The mean RTs and error percentages from the subject analyses are presented in Table 3. The design was the same as that in Experiment 1.

The ANOVA on the RTs showed that responses to word targets were faster when they were preceded by a related prime than when they were preceded by an unrelated prime, $F_1(1, 20) = 7.08, p < .02, F_2(1, 116) = 10.25, p < .001$. This priming effect was similar in size for word primes and transposed-letter nonword primes (15 vs. 16 ms, respectively), as deduced by the lack of interaction between relatedness and type of prime (both $F_s < 1$). The effect of type of prime was not significant (both $F_s < 1$).

The ANOVA for the error data revealed only a significant interaction between relatedness and type of prime, $F_1(1, 20) = 7.10, p < .02, F_2(1, 116) = 6.15, p < .02$. This interaction reflected a significant associative priming effect for target words preceded by a word prime, $F_1(1, 20) = 6.96, p < .02, F_2(1, 116) = 5.59, p < .025$, but not for target words preceded by a transposed-letter nonword prime, $F_1(1, 20) = 1.10, p > .30, F_2(1, 116) = 1.23, p > .26$.

The present experiment provides a successful replication to Experiment 1 with a sample of participants from another university: Associative priming can be obtained with transposed-letter nonword primes, and its magnitude does

not differ significantly from that of word primes (i.e., the pattern of data in Experiment 1 was not an empirical anomaly based on a unique participant population).

General discussion

The main finding of the present lexical decision experiments is that transposed-letter nonword primes (e.g., *jugde*) can sustain associative/semantic activation from their base words even when they are clearly visible—as deduced from the obtained associative priming effect with transposed-letter nonword primes obtained in Experiments 1–3. Importantly, these transposed-letter nonwords activate their base words to a large degree; the magnitude of associative priming for word primes and transposed-letter nonword primes in Experiments 1 and 3 was quite similar. The semantic activation from the transposed-letter nonwords is sustained in time; the magnitude of the associative priming effect for transposed-letter nonwords was similar when the prime exposure durations were 200 and 400 ms (Experiment 1).

Clearly, the timing corresponding to the encoding of identities of the individual letters and their positions is an issue that needs to be examined in greater depth by the recently developed input coding schemes in visual word recognition. None of the current implementations of the models of visual word recognition that employ a flexible coding scheme include an associative/semantic layer. Adding a semantic/associative layer does not ensure that the models will be able to account for the pattern of results presented in this article. In order for any of the current coding schemes to generate the correct level of semantic facilitation for the different types of nonword primes, the word-level activation generated by transposed-letter nonword primes has to be larger and more sustained than the one generated by replacement-letter nonword primes. To explore this issue, we conducted a simulation on the spatial

Table 3 Mean lexical decision times (in milliseconds) and percentages of errors (in parentheses) for word targets in Experiment 3

	Type of Prime		
	Related	Unrelated	Priming
Words	643 (3.5)	658 (6.4)	15 (2.9)
TL nonwords	645 (5.8)	661 (4.6)	16 (-1.3)

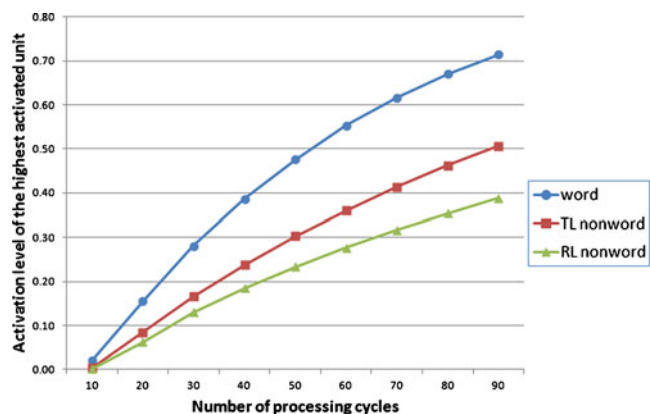


Fig. 1 Activation level of the most activated word unit in the spatial coding model (Davis, 2010), using the prime stimuli from the present experiments

coding model (Davis, 2010), using the default parameter values, and examined the degree of activation at the word level from words, transposed-letter nonwords, and replacement-letter nonwords. In these simulations, we individually presented the words, the transposed-letter nonwords, and the replacement-letter nonwords used in our experiments and collected the activation level of the most activated word unit as a function of the number of processing cycles. As can be seen in Fig. 1, transposed-letter nonwords yielded consistently higher levels of word-level activation than did replacement-letter nonwords. (Other models with flexible input coding schemes would also predict a similar pattern of data.) Thus, if the semantic layer were implemented, transposed-letter nonwords should generate more activation at the semantic level than the replacement-letter nonwords. It is important to note here that replacement-letter nonwords may generate some semantic activation from their base words: Bourassa and Besner (1998) demonstrated that there is a small associative/semantic priming effect when replacement-letter primes are presented briefly and masked (see also Perea & Lupker, 2003). What happens is that replacement-letter nonwords are less similar to their base words than are transposed-letter pseudowords, and hence, they are less likely to produce associative/semantic priming effects. Indeed, it is likely that, on a number of trials, replacement-letter nonword primes can activate the semantic information from their base words, as is actually suggested by the nonsignificant trend for replacement-letter nonword primes in Experiment 2.

The presence of a sustained semantic activation from transposed-letter nonwords like *jugde* is consistent with intuition: Even if we notice that there is a misspelling in *jugde*, subjective experience tells us that we can still access the meaning of the base word (i.e., *judge*). Converging evidence comes from a recent lexical decision experiment which collected evoked-response potentials (ERPs). Carreiras et al. (2007) reported that words and transposed-letter pseudowords (but not replacement-letter words) elicited quite similar ERP waves at the initial part of the N400 component, which is a component that has been associated with lexical-semantic processing (see Bentin, Kutas, & Hillyard, 1993). Taken together, the empirical evidence strongly suggests that transposed-letter nonwords can readily activate the lexical/semantic entries corresponding to their base words.

In sum, the present experiments have demonstrated that, unlike other types of nonwords, a transposed-letter nonword like *jugde* activates to some degree the meaning of COURT (via *judge*) even when the prime is clearly visible. This dissociation poses constraints to the computational models of visual word recognition that account for transposed-letter similarity effects.

Author note The research reported in this article has been partially supported by Grant PSI2008-04069/PSIC from the Spanish Ministry of Education and Science. We thank Derek Besner, Marc Brysbaert, and Keith Rayner for helpful criticism on an earlier version of the manuscript. We also thank Matt Lord and Robert Zimmerman for their support in collecting data.

References

- Bentin, S., Kutas, M., & Hillyard, S. A. (1993). Electrophysiological evidence for task effects on semantic priming in auditory word processing. *Psychophysiology*, 30, 161–169.
- Bourassa, D. C., & Besner, D. (1998). When do nonwords activate semantics? Implications for models in visual word recognition. *Memory & Cognition*, 26, 61–74.
- Bruner, J. S., & O'Dowd, D. (1958). A note on the informativeness of parts of words. *Language and Speech*, 1, 98–101.
- 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.
- Davis, C. J. (2010). The spatial coding model of visual word identification. *Psychological Review*, 117, 713–758.
- Drieghe, D., & Brysbaert, M. (2002). Strategic effects in associative priming with words, homophones, and pseudohomophones. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 28, 951–961.
- Duñabeitia, J. A., Carreiras, M., & Perea, M. (2008). Are coffee and toffee served in a cup? Ortho-phonologically mediated associative priming. *Quarterly Journal of Experimental Psychology*, 61, 1861–1872.
- Duñabeitia, J. A., Perea, M., & Carreiras, M. (2009). There is no clam with coats in the calm coast: Delimiting the transposed-letter priming effect. *Quarterly Journal of Experimental Psychology*, 62, 1930–1947.
- Forster, K. I., & Forster, J. C. (2003). DMDX: A Windows display program with millisecond accuracy. *Behavior Research Methods, Instruments, & Computers*, 35, 116–124.
- Gomez, P., Ratcliff, R., & Perea, M. (2008). The overlap model: A model of letter position coding. *Psychological Review*, 115, 577–601.
- Grainger, J., & van Heuven, W. (2003). Modeling letter position coding in printed word perception. In P. Bonin (Ed.), *The mental lexicon* (pp. 1–24). New York: Nova Science.
- Lesch, M. F., & Pollatsek, A. (1993). Automatic access of semantic information by phonological codes in visual word recognition. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 19, 285–294.
- Lukatela, G., & Turvey, M. T. (1994). Visual access is initially phonological: 1. Evidence from associative priming by words, homophones, and pseudohomophones. *Journal of Experimental Psychology: General*, 123, 107–128.
- McClelland, J. L., & Rumelhart, D. E. (1981). An interactive activation model of context effects in letter perception: Part 1. An account of basic findings. *Psychological Review*, 88, 375–407.
- O'Connor, R. E., & Forster, K. I. (1981). Criterion bias and search sequence bias in word recognition. *Memory & Cognition*, 9, 78–92.
- Paap, K. R., Newsome, S. L., McDonald, J. E., & Schvaneveldt, R. W. (1982). An activation-verification model for letter and word recognition: The word-superiority effect. *Psychological Review*, 89, 573–594.
- Perea, M., & Fraga, I. (2006). Transposed-letter and laterality effects in lexical decision. *Brain and Language*, 97, 102–109.

- Perea, M., & Lupker, S. J. (2003). Does *jugde* activate COURT? Transposed-letter confusability effects in masked associative priming. *Memory & Cognition*, 31, 829–841.
- Perea, M., & Rosa, E. (2002). Does the proportion of associatively related pairs modulate the associative priming effect at very brief stimulus-onset asynchronies? *Acta Psychologica*, 110, 103–124.
- Perea, M., Rosa, E., & Gómez, C. (2005). The frequency effect for pseudo-words in the lexical decision task. *Perception & Psychophysics*, 67, 301–314.
- Pollatsek, A., & Well, A. (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.
- Rayner, K., White, S. J., Johnson, R. L., & Liversedge, S. P. (2006). Raeding wrods with jubmled lettres: There is a cost. *Psychological Science*, 17, 192–193.
- 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.
- Yap, M. J., Tse, C.-S., & Balota, D. A. (2009). Individual differences in the joint effects of semantic priming and word frequency: The role of lexical integrity. *Journal of Memory and Language*, 61, 303–325.