

Can *CANISO* activate *CASINO*? Transposed-letter similarity effects with nonadjacent letter positions[☆]

Manuel Perea^{a,*} and Stephen J. Lupker^b

^a *Departament de Metodologia, Universitat de València, Facultat de Psicologia, Av. Blasco Ibáñez, 21, 46010 València, Spain*

^b *Department of Psychology, University of Western Ontario, London, Ont., Canada N6A 5C2*

Received 6 February 2004; revision received 14 May 2004

Available online 17 June 2004

Abstract

Nonwords created by transposing two *adjacent* letters (i.e., transposed-letter (TL) nonwords like *jugde*) are very effective at activating the lexical representation of their base words. This fact poses problems for most computational models of word recognition (e.g., the interactive-activation model and its extensions), which assume that exact letter positions are rapidly coded during the word recognition process. To examine the scope of TL similarity effects further, we asked whether TL similarity effects occur for nonwords created by exchanging two *nonadjacent* letters (e.g., *caniso-CASINO*) in three masked form priming experiments using the lexical decision task. The two nonadjacent transposed letters were consonants in Experiment 1 (e.g., *caniso-CASINO*), vowels in Experiment 2 (*anamil-ANIMAL*) and both consonants and vowels in Experiment 3. Results showed that nonadjacent TL primes produce priming effects (in comparison to orthographic controls, e.g., *caviro-CASINO*), however, only when the transposed letters are consonants. In a final experiment we examined latencies for nonwords created by nonadjacent transpositions of consonants versus vowels in a lexical decision task. Both types of nonwords produced longer latencies than matched controls, with consonant TL nonwords being more difficult than vowel TL nonwords. The implications of these findings for models having “position-specific” coding schemes as well as for models proposing alternative coding schemes are discussed.

© 2004 Elsevier Inc. All rights reserved.

Keywords: Coding schemes; Lexical decision; Orthographic similarity; Transposed letters; Vowel/consonant processing

Introduction

One key issue for models of visual word recognition is how the ordering of letters in a word is encoded

[☆] This research was supported by Natural Sciences and Engineering Research Council (NSERC) of Canada Grant A6333 to Stephen J. Lupker and a grant from the Spanish Ministry of Science and Technology (BSO2002-03286) to Manuel Perea. We thank Ken Forster, Carol Whitney, Iris Berent, and three anonymous reviewers for helpful criticism on an earlier draft. We also thank Edurne Goikoetxea for selecting the stimuli and conducting Experiments 3 and 4.

* Corresponding author. Fax: +1-34-96-3864697.

E-mail addresses: mperea@uv.es (M. Perea), lupker@uwo.ca (S.J. Lupker).

within that word's orthographic representation. Most current computational models of word recognition simply assume that the positions of the letters are established very early in processing, well before the identities of the letters are known (“position-specific” coding schemes; e.g., the interactive-activation (IA) model, Rumelhart & McClelland, 1982, and the models deriving from its architecture, the dual-route cascaded model, Coltheart, Rastle, Perry, Ziegler, & Langdon, 2001; and the multiple read-out model, Grainger & Jacobs, 1996). Thus, in these models, a nonword created by transposing two adjacent letters (e.g., *JUGDE*) would be no more similar to its base word (*JUDGE*) than a nonword created by simply replacing those letters (*JUNPE*).

This type of coding scheme does not, however, appear to fit with the available data. For example, a number of experiments have shown that people actually have more difficulty accurately perceiving letter order information than letter identity information when a random sequence of letters is briefly presented (e.g., Mewhort, Campbell, Marchetti, & Campbell, 1981; Ratcliff, 1981). Results have also shown that adjacent transposed-letter (TL) nonwords (e.g., *JUGDE*) have a strong tendency to be misperceived as words in a lexical decision task, a tendency that is even stronger than that for replacement-letter nonwords (*JUNGE*) (see Chambers, 1979; O'Connor & Forster, 1981). Further, in masked priming experiments (Forster & Davis, 1984), TL nonword primes not only produce form-priming effects relative to an orthographic control (e.g., *jugde-JUDGE* vs. *jupte-JUDGE*; Perea & Lupker, 2003b; see also Andrews, 1996; Forster, Davis, Schoknecht, & Carter, 1987; Peressotti & Grainger, 1999; Schoonbaert & Grainger, 2004), but also associative-priming effects (e.g., *jugde-COURT* vs. *ocaen-COURT*; Perea & Lupker, 2003a).

The presence of these “TL similarity effects” poses a challenge for word recognition models that use a position-specific coding scheme, that is, those models in which letters are assumed to be immediately assigned to their correct positions in the letter string. At the very least, addressing this challenge would require the models to drop this immediate assignment assumption and, instead, assume that letter positions often take more time to encode than letter identities. In addition, one would also need to assume that a letter in position N produces some activation of the representation of that same letter at positions $N - 1$ and $N + 1$ (Rumelhart & McClelland, 1982; see also Andrews, 1996; Peressotti & Grainger, 1995). Indeed, Rumelhart and McClelland (1982) acknowledged that there might be a problem with the coding scheme in their model and suggested that:

information presented in one location might activate detectors in a range of locations rather than simply in one fixed position. Perhaps there is a region of uncertainty associated with each feature and with each letter. If so a given feature in a given input position would tend to activate units for that feature in positions surrounding the actual appropriate position. As a result, partial activation of letters from nearby positions would arise in a particular position along with the activation for the letter actually presented. (p. 89)

Similarly, Peressotti and Grainger (1995) indicated that there could be “some form of cross-talk between neighboring letter positions” (p. 886), and Andrews (1996) indicated that “a letter in position n yields some activation of the same letter in positions $n - 1$ and $n + 1$ ” (p. 797).

With this reformulated coding scheme, the TL nonword *JUGDE* would be expected to activate the lexical

entry corresponding to its base word (*JUDGE*) somewhat more than the two-different letter nonword *JUNPE*, allowing the models to account for these types of effects. Simulation work would, of course, be necessary in order to verify that integrating these ideas into the letter coding schemes in the current models would allow those models to: (a) successfully capture TL similarity effects, and (b) maintain their ability to account for the other effects that they currently are able to account for (i.e., consider the discussion below concerning Davis's, 1999, analysis). In principle, however, the presence of TL similarity effects based on adjacent letters may not present an insurmountable challenge for a modified IA model.

In the present paper, we wished to examine the scope of TL similarity effects. Specifically, using the masked priming technique (Experiments 1–3), we asked whether TL similarity effects occur for TL nonwords created by transposing two *nonadjacent* letters (e.g., *caniso-CASINO*). In all cases, these effects were evaluated relative to the appropriate orthographic controls (i.e., replacement-letter nonwords as primes, as in *caviro-CASINO*). In an effort to obtain additional evidence on this issue, we also asked whether the TL nonwords created by transposing two *nonadjacent* letters are more competitive (in terms of the number of false positives and longer latencies) than their corresponding orthographic controls in a lexical decision task (Experiment 4).

The presence of TL priming effects when the transposed letters are not adjacent would pose a substantially greater problem for word recognition models using “position-specific” coding schemes. That is, these effects would require the assumption that a letter in position N activates its representation across letter positions $N - 2$ to $N + 2$. As described by Davis (1999), incorporating this assumption would seriously harm these models' ability to recognize highly familiar inputs. The reason is that, in the IA architecture, the mechanism responsible for resolving competition between candidate letters in a given letter position is bottom-up inhibition between the feature and letter levels. If a given letter position were receiving activation from features from up to five different letters (i.e., letters in positions $N - 2$ to $N + 2$) it would be almost impossible for the inhibitory process to function effectively. Thus, the existence of nonadjacent TL priming effects would strongly suggest that IA-based models would be best served by incorporating a different type of coding scheme.

The search for a new coding scheme

The existence of TL similarity effects, even when the transposed letters are not adjacent, is, in fact, a natural consequence of the letter coding schemes in two recently proposed computational models of the letter coding process: the SOLAR model (Davis, 1999) and the

SERIOR model (Whitney, 2001).¹ The SOLAR model uses a spatial coding scheme in which letter codes are position-independent, so that the nonadjacent TL nonword *CANISO* and its base word, *CASINO*, share the same set of letter nodes. The order of the letters is coded by the relative activity of the set of letter nodes. Thus, *CANISO* and *CASINO* would be coded differently because they would produce different activation patterns across the letter nodes they share (e.g., in the word *CASINO*, the letter node corresponding to *C* is the one associated with the highest activation value, the letter node corresponding to the letter *A* is associated with a slightly smaller activation value, etc.). The SERIOR model (Whitney, 2001) uses a “letter-tagging” coding scheme, in which each letter is marked for the ordinal position in which it occurs within a letter string. For instance, the word *CASINO* would be represented by *C-1*, *A-2*, *S-3*, *I-4*, *N-5*, and *O-6* with the relevant letter nodes then receiving differential levels of activation as a function of position. This letter-tagging scheme is accompanied by the activation of bigram nodes—ordered pairs of letters—so that *CASINO* would be represented by the following bigram nodes: *CA*, *AS*, *SI*, *IN*, *NO*, *CS*, *CI*, *CN*, *CO*, *AS*, *AI*, *AN*, *AO*, *SI*, *SN*, *SO*, *IN*, and *IO*. The nonadjacent TL nonword *CANISO* would then share 13 bigram nodes with *CASINO* (*CA*, *AS*, *NO*, *CS*, *CI*, *CN*, *CO*, *AS*, *AI*, *AN*, *AO*, *SO*, and *IO*), whereas the two-letter different nonword *CAVIRO* would share only six bigram nodes with *CASINO* (*CA*, *CI*, *CO*, *AI*, *AO*, and *IO*).

For our purposes, the crucial point here is that, according to both the SOLAR and SERIOR models, nonwords created by transposing nonadjacent letters are highly similar to their base words. Further, although the precise similarity of *CASINO* and its nonadjacent TL

nonword *CANISO* depends on a variety of factors, the current parameter sets in the SOLAR and SERIOR models predict that a nonadjacent TL nonword like *CANISO* is more similar to *CASINO* than a two-letter different nonword like *CAVIRO*. More specifically, in terms of calculated similarity, the similarity match between *CASINO* and *CANISO* would be 1.00 in both models. For the SOLAR and SERIOR models, respectively, the similarity match to *CASINO* would be reduced to .83 or .88 for the adjacent TL neighbor *CAISNO*, to .75 or .83 for a one-letter different nonword like *CASIRO*, to .62 or .71 for the nonadjacent TL nonword *CANISO*, and to .54 or .49 for a two-letter different nonword like *CAVIRO* (an unrelated nonword like *NOMERA* results in a match value of .13 in the SOLAR model and a match value of .20 in the SERIOR model).² Thus, although the presence of TL similarity effects based on nonadjacent letter positions would pose considerable problems for position-specific coding schemes, their presence would actually support the predictions of these two recent models of letter coding.

In Experiments 1–3 we asked whether TL nonword primes created by transposing two nonadjacent letters do produce reliable form-priming effects relative to two-letter different nonwords (e.g., *caniso-CASINO* vs. *caviro-CASINO*). Because TL similarity effects are greater for nonwords created by transposing internal rather than external letters (Chambers, 1979; Perea & Lupker, 2003a, 2003b), the transposed letters were always the third and the fifth in six-letter words (or nonwords) in Experiments 1 and 2. (They were the third and fifth, or the fourth and the sixth in Experiment 3, in which the items were 7–10 letters long; e.g., *tragedia-TRAGEDIA*.) For comparison purposes, we also included a one-letter different condition (*casiro-CASINO*), as well as either an identity condition (*casino-CASINO*; Experiment 1a), or an unrelated condition (*numero-CASINO*, Experiments 1b and 2). The two nonadjacent transposed letters were consonants in Experiments 1a and 1b (e.g., *caniso-CASINO* vs. *caviro-CASINO*) and vowels in Experiment 2 (*anamil-ANIMAL* vs. *anamel-ANIMAL*). Experiment 3 was an attempt to replicate the results of Experiments 1 and 2 using a new set of items. To obtain converging evidence of TL similarity effects using another experimental technique, Experiment 4 employed a single-presentation lexical decision task in which the nonword targets were the masked primes used in Experiment 3.

The differential processing of vowels and consonants

The second issue being addressed in the present research concerns potential differences in the processing of vowels versus consonants. Recent research strongly

¹ Other coding schemes had been proposed in the literature, the more cited of which has been Mozer's (1987) BLIRNET model. Mozer (1987) used letter-cluster units that respond to local arrangements of letters in which the only location information retained consisted of the relative positions of letters within a cluster. The letter-cluster units respond to letter triples: either a sequence of three adjacent letters (e.g., *CAS* in the Spanish word *CASINO*), or two adjacent letters and one nearby letter, such as *CA_I* or *S_NO*, where the line indicates that any letter may appear in the corresponding position. For instance, presentation of *CASINO* should result in the activation of the following letter-cluster units: ***C*, ***_A*, **CA*, **_AS*, **C_S*, *CAS*, *C_SI*, *CA_I*, *ASI*, *A_IN*, *AS_N*, *SIN*, *S_NO*, *SI_O*, *INO*, *I_O**, *IN_**, *NO**, *N_**, and *O***. (The asterisk signifies a blank space and double-asterisks are used simply to keep all units in the form *xxx*, *xx_x*, or *x_xx*.) The word *CASINO* and the nonadjacent TL pseudoword *CANISO* only share 5 (out of 20) letter-cluster units. The word *CASINO* and the orthographic control *CAVIRO* also share five letter-cluster units. Thus, BLIRNET would predict that nonadjacent TL nonwords are no more similar to their base words than orthographic controls.

² We thank Colin Davis and Carol Whitney for providing us with the match scores.

suggests that vowels and consonants are processed differently (visual-word perception: Berent & Perfetti, 1995; Berent, Bouissa, & Tuller, 2001; Lee, Rayner, & Pollatsek, 2001, 2002; speech perception: Boatman, Hall, Goldstein, Lesser, & Gordon, 1997; neuropsychology: Caramazza, Chialant, Capasso, & Miceli, 2000; Caramazza & Miceli, 1990; Cubelli, 1991; Ferreres, López, Petracci, & China, 2000; cognitive modeling, Monaghan & Shillcock, 2003; language acquisition, Nespor, Peña, & Mehler, 2003; linguistics, Gafos, 1998; Goldsmith, 1990). The experimental tasks used in this research varied widely, from perceptually based tasks to production tasks (see Berent et al., 2001; Monaghan & Shillcock, 2003). The clear implication, as a number of these researchers have suggested, is that the consonant/vowel distinction is an essential one in that orthographic representations convey information concerning not only letter identity but also consonant/vowel status (e.g., Berent et al., 2001; Caramazza & Miceli, 1990; Tainturier & Caramazza, 1996).

Note, as well, that the dimension of consonant/vowel status also emerges as important in patterns of brain damage. That is, there are documented cases of brain damaged patients with a selective deficit with vowels (e.g., Caramazza et al., 2000; Cubelli, 1991) or with consonants (e.g., Caramazza et al., 2000; Kay & Hanley, 1994). Thus, it is even possible that consonants and vowels are processed by different neural mechanisms (Caramazza et al., 2000).

Finally, it is worth noting two additional facts particularly relevant to the present research. First, in normal speech in Spanish (the language used in the present experiments), individuals make substantially more pronunciation errors by transposing two consonants (14.2% of errors; e.g., the nonword *escanLaDosa* instead of the Spanish word *escanDaLosa*) than by transposing two vowels (1.9%; see Pérez, Palma, & Santiago, 2001). Second, a post hoc analysis of Experiment 3 of Perea and Lupker (2003b) revealed that transposition of two adjacent internal consonants (e.g., *mohter-MOTHER*) led to as much priming as that from identity primes, whereas the transposition of two adjacent internal vowels (*freind-FRIEND*) produced very little priming. Unfortunately, this comparison is not only post hoc, it also fails to control for any differences between within- and between-syllable transpositions. In the present experiments, due to the syllable structure of Spanish, all transpositions were between-syllable transpositions.

How the consonant/vowel distinction actually manifests itself during processing has been a matter of considerable theorizing in recent years. For instance, Berent et al. (2001) have proposed that printed words are represented in the internal lexicon in terms of a consonant/vowel skeletal structure, in which there are different slots for consonants and vowels. Caramazza et al. (2000) proposed that letters are classified as consonants or

vowels and this categorical distinction then plays a key role in processes such as the construction of syllables in speech production. Nespor et al. (2003) suggested that the rapid classification of letters as consonants or vowels allows a division of labor between their processing, with vowels being used to help the reader interpret grammar, whereas the role of consonants is to aid in accessing the internal lexicon. The important point here is that there is now considerable opinion that consonants and vowels are processed differently. Thus, although none of these theories would appear to make any specific predictions as to whether TL similarity effects would vary as a function of whether the transposed letters are vowels or consonants, any theories proposing consonant/vowel differences would be informed by the results of such a comparison.

For these reasons, it seemed important to examine whether TL similarity effects do differ as a function of whether the transposed letters are vowels or consonants. What should be explicitly noted is that because the orthographic representations in the SERIOL and SOLAR models do not convey consonant/vowel status, TL similarity effects with nonadjacent letters should be the same for consonant and vowel transpositions. Thus, if consonant/vowel differences are observed in the present experiments, at the very least, these models will need an additional mechanism in order to explain those differences.

Experiment 1

The prime types in Experiments 1a were: (1) identity (*casino-CASINO*), (2) one-letter replacement nonword (*casiro-CASINO*), (3) nonadjacent TL nonword (*caniso-CASINO*) and (4) two-letter replacement nonword (*caviro-CASINO*). As noted, both SERIOL and SOLAR predict that the latencies for these four conditions should increase monotonically between the first and fourth conditions. The only difference between Experiments 1a and 1b was that the identity condition in Experiment 1a was replaced by an unrelated nonword condition in Experiment 1b. Thus, Experiment 1b allowed us to better gauge the sizes of the form priming effects in the other three conditions (by comparing their latencies to that in the unrelated condition). It also allowed us the opportunity to replicate the main finding of Experiment 1a, the advantage of the nonadjacent TL prime condition over the two-letter replacement letter prime condition.

Method

Participants

Fifty-six students from the University of València received course credit for participating in the experiment

(28 in Experiment 1a and 28 in Experiment 1b). All of them either had normal or corrected-to-normal vision and were native speakers of Spanish.

Materials

The targets were 128 Spanish words of six letters (mean word frequency per one million words in the Alameda & Cuetos, 1995, count: 42, range: 2–418; mean Coltheart's *N*: 2.3, range: 0–11). The targets in Experiment 1a were presented in uppercase and were preceded by primes in lowercase that were: (1) the same as the target (identity condition), e.g., *casino-CASINO*, (2) the same except for the substitution of one internal letter (always the fifth letter; one-letter different condition), *casiro-CASINO*, (3) the same except for a transposition of the third and the fifth letters (nonadjacent TL condition), *caniso-CASINO*, and (4) the same except for the substitution of two internal letters (the third and the fifth letters; two-letter different condition), *caviro-CASINO*. Except in the identity condition, the primes were always nonwords. An additional set of 128 nonwords of six letters was included for the purposes of the lexical decision task (mean Coltheart's *N*: 1.0, range: 0–5). The manipulation of the nonword trials was the same as that for the word trials. All items had a CV.CV.CV syllabic structure. Four lists of materials were constructed so that each target appeared once in each list, but each time in a different priming condition. Different groups of participants were used for each list. The materials of Experiment 1b were the same as in Experiment 1a, except that the identity primes were replaced by unrelated nonword primes (e.g., *numero-CASINO*). The related pairs are given in the Appendix.

Procedure

Participants were tested in groups of four to eight in a quiet room. Presentation of the stimuli and recording of response times were controlled by Apple Macintosh Classic II microcomputers. The routines for controlling stimulus presentation and reaction time collection were obtained from Lane and Ashby (1987) and from Westall, Perkey, and Chute (1986), respectively. Reaction times were measured from target onset until the partic-

ipant's response. On each trial, a forward mask consisting of a row of six hash marks (#####) was presented for 500 ms in the center of the screen. Next, a centered lowercase prime was presented for 50 ms. Primes were immediately replaced by an uppercase target item, which remained on the screen until the response. Participants were instructed to press one of two buttons on the keyboard to indicate whether the uppercase letter string was a legitimate Spanish word or not ("ç" for yes and "z" for no). Participants were instructed to make this decision as quickly and as accurately as possible. Participants were not informed of the presence of lowercase items. Each participant received a different order of trials. Each participant received a total of 20 practice trials (with the same manipulation as in the experimental trials) prior to the 256 experimental trials. The whole session lasted approximately 15 min.

Results and discussion

Incorrect responses (6.7% of the data for word targets) and reaction times less than 250 ms or greater than 1500 ms (0.6% of the data for word targets) were excluded from the latency analysis. Although the critical contrast was the comparison between the nonadjacent TL condition and its control condition (i.e., the two-letter different priming condition), we also conducted *F* tests, based on both the subject (*F*1) and item (*F*2) means, for the following contrasts: one-letter different primes vs. two-letter different primes, one-letter different primes vs. identity primes (for Experiment 1a), and two-letter different primes vs. unrelated primes (for Experiment 1b). To extract the variance due to the error associated with the lists, List was included as a dummy variable in all comparisons. All significant effects had *p* values less than the .05 level. The mean response times and error percentages from the subject analysis are presented in Table 1.

Experiment 1a

Word data. Targets preceded by a nonadjacent TL prime were responded to 21 ms faster than the targets preceded by a two-letter different prime, $F(1, 24) =$

Table 1
Mean lexical decision times (in ms) and percentage of errors (in parentheses) for word and nonword targets in Experiment 1

	Type of prime				
	Identity	One-letter different	NonAdj. TL	Two-letter different	Unrelated
<i>Experiment 1a</i>					
Word trials	598 (4.5)	615 (6.9)	628 (7.8)	649 (8.6)	
Nonword trials	695 (6.3)	691 (3.9)	694 (5.0)	696 (3.3)	
<i>Experiment 1b</i>					
Word trials		665 (6.1)	679 (6.4)	696 (6.4)	703 (6.7)
Nonword trials		822 (5.4)	814 (7.1)	813 (6.4)	832 (6.0)

13.00; $F_2(1, 124) = 10.02$. In addition, targets preceded by a one-letter different prime were responded to 34 ms faster than the targets preceded by a two-letter different prime, $F_1(1, 24) = 42.18$; $F_2(1, 124) = 35.58$, and targets preceded by an identity prime were responded to 13 ms faster than the targets preceded by a one-letter different prime, $F_1(1, 24) = 14.14$; $F_2(1, 124) = 18.76$. None of the contrasts on the error data were statistically significant (all $ps > .10$).

Nonword data. In the latency analyses, there were virtually no differences across the different priming conditions. With respect to the error data, participants committed 1.7% more errors to nonwords preceded by a nonadjacent TL prime than to nonwords preceded by a two-letter different prime (5.0 vs. 3.3%, respectively), $F_1(1, 24) = 7.42$; $F_2(1, 124) = 3.10$, $p = .08$. The other two contrasts did not approach significance (all $ps > .10$).

Experiment 1b

Word data. As in Experiment 1a, targets preceded by a nonadjacent TL prime were responded to 17 ms faster than the targets preceded by a two-letter different prime, $F_1(1, 24) = 7.47$; $F_2(1, 124) = 5.21$, and targets preceded by a one-letter different prime were responded to 31 ms faster than the targets preceded by a two-letter different prime, $F_1(1, 24) = 42.18$; $F_2(1, 124) = 20.11$. Finally, the 7-ms difference between targets preceded by a two-letter different prime and the targets preceded by an unrelated prime was not significant, both $ps > .10$. The error analyses did not reveal any significant effects (all $ps > .10$).

Nonword data. There were virtually no differences across the different priming conditions in either the latency or error data.

The results were straightforward. There was a sizable priming effect (17–21 ms) from nonadjacent TL nonword primes relative to the appropriate orthographic

control condition (i.e., the two-letter different condition) in both Experiments 1a and 1b (also see Fig. 1). Thus, it appears that nonadjacent TL nonwords do activate, to a greater degree than two-letter different nonwords, the lexical representation of their base words. Note also that the pattern of priming effects across the various conditions is consistent with the predictions of the SERIOL and SOLAR models using their default parameter settings (see Introduction). The only possible exception to this is the comparison between the unrelated word primes and the two-letter different primes. Although two-letter different primes are much more similar to their base words than unrelated primes according to both models, there was no significant latency difference between the two conditions.

Experiment 2

The prime–target conditions in Experiment 2 were the same as in Experiment 1b, except that the transposed letters were vowels (e.g., *anamil-ANIMAL* vs. *anemol-ANIMAL*) instead of consonants.

Method

Participants

Twenty-four students from the same population as in Experiment 1 participated in this experiment.

Materials

The word targets were 128 Spanish words of six letters (mean word frequency per one million words in the Alameda & Cuetos, 1995, count: 28, range: 1–379; mean Coltheart's N : 1.4, range: 0–8) and 128 nonwords of six letters (mean Coltheart's N : 0.6, range: 0–5). All word and nonword targets had vowels in positions three and

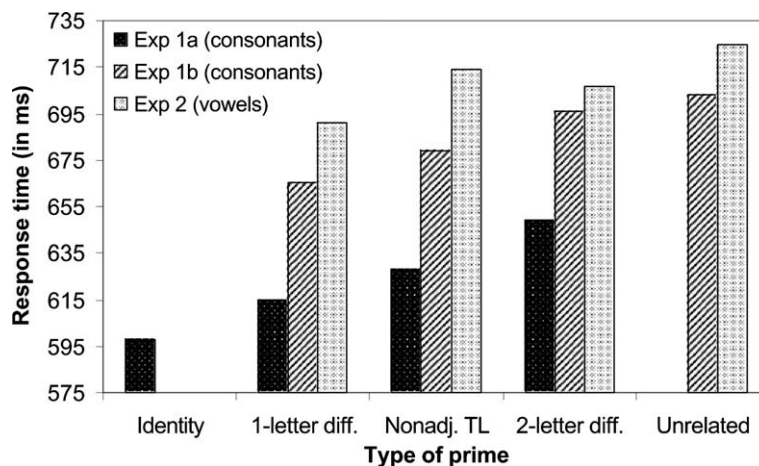


Fig. 1. Response times for word targets in the different prime–target relationships in Experiments 1 and 2.

five; e.g., the word *ANIMAL*, or the nonword *ECUDAR*. The prime–target conditions were the same as in Experiment 1b: (1) the one-letter different condition, e.g., *animol-ANIMAL*, (2) the nonadjacent TL condition, *anamil-ANIMAL*, (3) the two-letter different condition, *anomal-ANIMAL*, and (4) the unrelated condition, *agapón-ANIMAL*. Primes were always nonwords. The manipulation for the nonword trials was the same as that for the word trials. Four lists of materials were constructed so that each target appeared once in each list, but each time in a different priming condition. Different groups of participants were used for each list. The related pairs are given in the Appendix.

Procedure

The procedure was the same as in Experiment 1.

Results and discussion

Incorrect responses (5.9% of the data for word targets) and reaction times less than 250 ms or greater than 1500 ms (1.9% of the data for word targets) were excluded from the latency analysis. As in Experiment 1, the critical contrast was the comparison between the nonadjacent TL condition and its orthographic control condition (i.e., the two-letter different condition). The mean response times and error percentages from the subject analysis are presented in Table 2.

Word data. Unlike in Experiment 1, targets preceded by a nonadjacent TL prime were responded to 9 ms slower (rather than faster) than the targets preceded by a two-letter different prime. This difference was not significant (both $ps > .20$). In addition, targets preceded by a one-letter different prime were responded to 15 ms faster than the targets preceded by a two-letter different prime, $F1(1, 20) = 6.02$; $F2(1, 124) = 6.85$, and targets preceded by a two-letter different prime were responded to 20 ms faster than the targets preceded by an unrelated prime, $F1(1, 20) = 9.54$; $F2(1, 124) = 4.50$. Finally, the 11-ms difference between the nonadjacent TL priming condition and the unrelated priming condition approached significance in the subject analysis, $F1(1, 20) = 4.09$, $p = .057$; $F2(1, 124) = 1.20$. The error analyses did not reveal any significant effects.

Nonword data. There were virtually no differences across the different priming conditions in either the latency or error data.

Although the results of Experiment 2 parallel those of Experiment 1 in some ways, the central finding is that priming effects from nonadjacent TL nonwords do not seem to occur when the TL nonwords are created by transposing two vowels (see Fig. 1).

Experiment 3

The main goal of Experiment 3 was to establish the generality of the priming effects obtained in Experiments 1 and 2, with a new set of (7–10 letter long) words. For reasons of design efficiency, we focused on the transposed-letter condition and its corresponding orthographic control condition, both for consonants (e.g., *tragedia-TRAGEDIA* vs. *trabepia-TRAGEDIA*) and for vowels (*absuloto-ABSOLUTO* vs. *abselito-ABSOLUTO*). The SOA was the same as in Experiments 1 and 2 (50 ms).

Method

Participants

Sixty-two students from the Universidad de Deusto took part in the experiment. All of them either had normal or corrected-to-normal vision and were native speakers of Spanish.

Materials

The targets were 80 Spanish words that were 7–10 letters long. Forty of these words (mean word frequency per one million words in the Alameda & Cuetos, 1995, count: 69, range: 28–210; mean Coltheart's N : 0.47, range: 0–2) were presented in uppercase and were preceded by primes in lowercase that were: (1) the same except for a transposition of two internal consonants (either letter positions 3–5 or 4–6; nonadjacent TL-consonant condition), *tragedia-TRAGEDIA*, or (2) the same except for the substitution of the corresponding internal consonants, *trabepia-TRAGEDIA*. The remaining forty of these words (mean word frequency per one million words in the Alameda & Cuetos, 1995; count: 68, range: 31–143; mean Coltheart's N : 0.37, range: 0–2) were also presented in uppercase and were preceded by primes in lowercase that were: (1) the same except for a transposition of two internal vowels (either letter positions 3–5 or 4–6; nonadjacent TL-vowel condition), *absuloto-ABSOLUTO*, or

Table 2

Mean lexical decision times (in ms) and percentage of errors (in parentheses) for word and nonword targets in Experiment 2

	Type of prime			
	One-letter different	Nonadj TL	Two-letter different	Unrelated
Word trials	691 (4.9)	714 (6.2)	707 (6.0)	725 (6.5)
Nonword trials	819 (4.3)	818 (5.5)	814 (6.3)	820 (4.2)

(2) the same except for the substitution of the corresponding internal vowels, *abselito-ABSOLUTO*. Primes were always nonwords. An additional set of 80 target nonwords that were 7–10 letters long was included for the purposes of the lexical decision task. The manipulation of the nonword trials was the same as that for the word trials. Two lists of materials were constructed so that each target appeared once in each list, but each time in a different priming condition. Different groups of participants were used for each list. The pairs are given in the Appendix.

Procedure

Participants were tested individually in a quiet room. Presentation and timing of stimuli were controlled by DMDX software (Forster & Forster, 2003) on a PC computer. Reaction times were measured from target onset until the participant's response. On each trial, a forward mask consisting of a row of ten hash marks (#####) was presented for 500 ms in the center of the screen. Next, a centered lowercase prime was presented for 50 ms. Primes were immediately replaced by an uppercase target item, which remained on the screen until the response. Participants were instructed to press one of two buttons on the keyboard to indicate whether the uppercase letter string was a legitimate Spanish word or not ("m" for yes and "z" for no). Participants were instructed to make this decision as quickly and as accurately as possible. Participants were not informed of the presence of lowercase items. Each participant received a different order of trials. Each participant received a total of 24 practice trials (with the same manipulation as in the experimental trials) prior to the 160 experimental trials. The whole session lasted approximately 12 min.

Results and discussion

Incorrect responses (3.8% of the data for word targets) and reaction times less than 250 ms or greater than 1500 ms (0.2% of the data for word targets) were excluded from the latency analysis. The mean response times and error percentages from the subject analysis are presented in Table 3. ANOVAs based on the subject and item mean correct response latencies and error rates were conducted based on a 2 (Relatedness: transposition, control) $\times$ 2 (Type of transposition/replacement: consonants, vowels) $\times$ 2 (List: list 1, list 2) design.

Word data. Targets preceded by a nonadjacent TL prime were responded to 12 ms faster than the targets preceded by a two-letter different prime, $F(1, 60) = 21.42$; $F(1, 76) = 18.55$. Response times to targets that involved the transposition/replacement of two consonants did not differ from the response times to targets that involved the transposition/replacement of two vowels, both F s < 1 . More important, the interaction of

Table 3

Mean lexical decision times (in ms) and percentage of errors (in parentheses) for word and nonword targets in Experiment 3

	Type of prime		
	Nonadj TL	Two-letter different	Priming
<i>Word trials</i>			
Consonants	573 (3.6)	591 (5.3)	18 (1.7)
Vowels	581 (3.5)	587 (2.9)	6 (–0.6)
<i>Nonword trials</i>			
Consonants	697 (8.5)	700 (7.3)	3 (–1.2)
Vowels	708 (7.5)	705 (8.6)	–3 (1.1)

the two factors was significant, $F(1, 60) = 4.37$; $F(1, 76) = 3.98$: there was a significant 18-ms relatedness effect for the transposed consonants, $F(1, 60) = 32.29$; $F(1, 38) = 29.35$, whereas there was a small (6 ms) nonsignificant relatedness effect for the transposed vowels, $F(1, 60) = 2.19$, $p > .10$; $F(1, 38) = 2.02$, $p > .10$.

The ANOVA on the error data showed a similar pattern: The interaction of the two factors was also significant, $F(1, 60) = 4.94$; $F(1, 76) = 4.88$: there was a 1.7% relatedness effect for the transposed consonants, $F(1, 60) = 6.52$; $F(1, 38) = 4.71$, whereas there was a –0.6% nonsignificant relatedness effect for the transposed vowels, both F s < 1 .

Nonword data. The ANOVAs on the nonword data did not reveal any significant effects (all p s $> .10$).

The results were again clear-cut. When the transposed letters were consonants (*tradegia-TRAGEDIA*), there was a sizable priming effect (18 ms) from nonadjacent TL nonword primes relative to the orthographic control condition (i.e., the two-letter different condition; *trabepia-TRAGEDIA*). When the transposed letters were vowels, there was only a nonsignificant 6 ms effect.

Experiment 4

The results of the first three experiments have clearly shown that nonadjacent TL-consonant nonwords activate, to a greater degree than two-letter different nonwords, the lexical representation of their base words. Further, they demonstrate that this effect seems to occur only for consonants (Experiments 1a, 1b, and 3) and not for vowels (Experiments 2 and 3). The goal of Experiment 4 was to obtain converging evidence on the role of nonadjacent TL consonants vs. vowels using another experimental technique: a single-presentation lexical decision task. We used the masked primes of Experiment 3 as the nonword targets. A set of word targets was selected for the purposes of the lexical decision task.

It is a well established finding that "wordlike" nonwords (e.g., one-letter different nonwords) produce

slower “no” responses and more errors than “non-wordlike” nonwords (e.g., those with no “orthographic” neighbors) in lexical decision tasks (e.g., Coltheart, Davelaar, Jonasson, & Besner, 1977; Forster & Shen, 1996; Perea & Rosa, 2000; Sears, Hino, & Lupker, 1995; Siakaluk, Sears, & Lupker, 2002). The explanation is that the former partially activate the lexical representations of their word neighbors. As a result, additional time is needed for the activation levels to settle and for the participant to realize that no word unit is being activated over threshold. Similarly, if nonadjacent TL nonwords (e.g., *TRADEGIA*) activate the lexical representation of their corresponding base words (*TRAGEDIA*) to a higher degree than orthographic controls (*TRABEPIA*), one would expect a higher rate of “word” responses and longer latencies for the nonadjacent TL nonwords than for the orthographic controls in the present experiment. Finally, based on the results of the previous experiments, we would expect this effect to exist for consonant transpositions but not necessarily for vowel transpositions.

One final point should be noted. There is good evidence for syllable-frequency effects in single-word presentation lexical decision tasks in Spanish (e.g., Carreiras & Perea, 2004; Perea & Carreiras, 1998). These effects, however, are restricted to the frequency of the initial syllable. Therefore, in order to avoid any uncontrolled effects of initial syllable frequency, all the experimental nonwords in Experiment 4 maintained the initial syllable of their base words. That is, the transposed/replacement letters all came from later syllables, meaning that the experimental and control nonwords had the same initial syllables.

Method

Participants

Twenty-six students from the Universidad de Deusto took part in the experiment. All of them either had normal or corrected-to-normal vision and were native speakers of Spanish.

Materials

The 80 word targets from Experiment 3 were used as the base words for the four nonword conditions (nonadjacent TL-consonant nonword and its corresponding control—e.g., *TRADEGIA* and *TRABEPIA*, nonadjacent TL-vowel nonword and its corresponding control—e.g., *ABSULOTO* and *ABSELITO*). That is, the nonword targets in the present experiment had been the nonword primes in the previous experiment. The TL-consonant nonwords and their orthographic controls both had, on average, 0.075 word neighbors (range 0–1) (note that all these neighbors were always very-low-frequency words, with a frequency no higher than 3 per million). The mean token trigram frequencies for the

critical trigrams in the transposed and the replacement-letter conditions (e.g., *DEG* and *BEP* in *TRAGEDIA* and *TRAPEDIA*) were 255 and 226 per million ($t < 1$), respectively. All these trigrams occurred in Spanish words. The mean token bigram frequencies for the critical bigrams in the two conditions (e.g., *AD*, *DE*, *EG*, *GI* in *TRADEGIA* vs. *AP*, *PE*, *ED*, *DI* in *TRAPEDIA*) were comparable (22,780 vs. 20,560 per million, for the transposed and replacement-letter conditions, respectively; $p > .15$). The TL-vowel nonwords and their orthographic controls had no orthographic neighbors. The mean token trigram frequencies for the critical trigrams in the transposed- and replacement-letter conditions (e.g., *ULO* and *ELI* in *ABSULOTO* and *ABSELITO*) were 632 and 583 per million ($t < 1$), respectively. All these trigrams occurred in Spanish words. (Note that the trigrams were more frequent for VCV than for CVC transpositions/replacements because of the obvious fact that there are only five vowel letters.) The mean token bigram frequencies of the critical bigrams in the two conditions (e.g., *SU*, *UL*, *LO*, *OT* in *ABSULOTO* vs. *SE*, *EL*, *LI*, *IT* in *ABSELITO*) were comparable (24,996 vs. 26,230 per million in the transposed- and replacement-letter conditions, respectively; $p > .15$). As noted, in all cases, the first syllable of the base word remained unchanged.

Two lists of materials were constructed so that if the TL-consonant nonword *TRADEGIA* appeared in one list, its orthographic control (*TRABEPIA*) would appear in the other list. Different groups of participants were used for each list. An additional set of 80 words that were 7–10 letters long (mean frequency per million words: 38, range: 1–206) was included for the purposes of the lexical decision task.

Procedure

Participants were tested individually in a quiet room. Presentation and timing of stimuli were controlled by DMDX software (Forster & Forster, 2003) on a PC computer. On each trial, a centered uppercase target item remained on the screen until response. Participants were instructed to press one of two buttons on the keyboard to indicate whether the letter string was a legitimate Spanish word or not (“m” for yes and “z” for no). Participants were instructed to make this decision as quickly and as accurately as possible. Each participant received a different order of trials. Each participant received a total of 24 practice trials (with the same manipulation as in the experimental trials) prior to the 160 experimental trials. The whole session lasted approximately 12 min.

Results and discussion

Incorrect responses (19.5% of the data for nonword targets, 5.3% for the word targets) and reaction times

Table 4
Mean lexical decision times (in ms) and percentage of errors (in parentheses) for nonword targets in Experiment 4

	Type of nonword		
	Nonadj TL	Two-letter different	Difference
Consonants	943 (43.5)	869 (4.6)	74 (38.9)
Vowels	915 (24.4)	853 (5.4)	62 (19.0)

Note. The mean correct RT for word trials was 752 ms and the error rate was 5.3%.

less than 250 ms or greater than 1500 ms (6.2% of the data for nonword targets) were excluded from the latency analysis. The mean response times and error percentages from the subject analysis are presented in Table 4. ANOVAs, based on both subject and item mean response latencies and error rates to nonword targets, were conducted based on a 2 (Type of nonword: transposition, control) $\times$ 2 (Type of transposition/replacement: consonants, vowels) $\times$ 2 (List: list 1, list 2) design.

Nonword targets created by transposing two non-adjacent letters were responded to 68 ms slower than nonwords created by replacing those two letters, $F(1, 24) = 34.86$; $F(1, 76) = 44.48$. Nonwords created by transposing/replacing two consonants had slower latencies than nonwords created by transposing/replacing two vowels, $F(1, 24) = 4.60$; $F(1, 76) = 3.97$. The difference between the transposed-letter nonwords and the replacement-letter nonwords was only slightly greater for the transpositions involving consonants than for the transpositions involving vowels (74 vs. 62 ms). Hence, the interaction between the two factors was not significant, both $ps > .10$.³

The ANOVA on the error data showed that there were significantly fewer errors to replacement-letter nonwords than to transposed-letter nonwords, $F(1, 24) = 116.21$; $F(1, 76) = 207.63$. There were also more errors to nonwords created by transposing/replacing two consonants than to nonwords created by

transposing/replacing two vowels, $F(1, 24) = 23.06$; $F(1, 76) = 15.24$. The interaction of the two factors was also significant, $F(1, 24) = 54.22$; $F(1, 76) = 24.31$: the transposition-letter effect was substantially larger for the nonwords created by transposing two consonants (38.9%) than for the nonwords created by transposing two vowels (19.9%). These two simple main effects were both significant: for consonants, $F(1, 24) = 168.15$; $F(1, 38) = 167.23$, for vowels, $F(1, 24) = 40.00$; $F(1, 38) = 50.95$.

Consistent with the previous experiments, transposed-letter nonwords created by transposing two non-adjacent consonants seem to activate their base word to a considerable degree. The effect was quite dramatic in terms of error rates (43.5 vs. 4.6%) for the transposed-letter condition and its orthographic control, respectively. One difference between the present experiment and the previous experiment is that this time, there were significant effects for transposing two nonadjacent vowels. (With these same materials, using the masked priming technique, there was only a nonsignificant 6-ms priming effect in Experiment 3.) In fact, only in the error data was there clear evidence that the effect of transposing two vowels was smaller than the effect of transposing two consonants.⁴ It appears, therefore, that vowel transposition nonwords are perceptually similar to their base words, they are simply less similar than consonant transposition nonwords.

Finally, it is important to note that the very high error rates for the TL nonwords do not reflect a very lenient decision criterion for “word” responses, but rather they reflect the high degree of similarity between the TL nonwords and their corresponding base words. Experiment 4 has been replicated with a different set of items in another lab in Spain, and the error rates for the TL-vowel nonwords and the TL-consonant nonwords were virtually the same (41 vs. 22%, respectively) as those reported here (43 vs. 24 %). Indeed, for any skilled reader of Spanish, the TL nonword *DEYASUNO* seems to activate its base word (*DESAYUNO*) to a large degree, and it is rather difficult to process/pronounce correctly a TL nonword such as *DEYASUNO* under time pressure.

³ The lack of an interaction between Type of nonwords and Type of transposition/replacement was not due to the use of a 1500 ms cutoff. With a more conservative cutoff, the critical interaction was also not significant (e.g., with a 2000 ms cutoff, the TL effect was 94 vs. 79 ms, for the consonant TL nonwords and for the vowel TL nonwords, respectively). Nonetheless, there were some signs of a greater TL effect for consonant nonwords than for vowel nonwords when examining the RT distributions (using *all* the correct responses): As in the mean RT analyses, the TL effect was only slightly greater for consonant nonwords than for vowel nonwords in the bulk of the RT distributions (.5 quantile: 101 vs. 89 ms for the consonant and the vowel TL nonwords, respectively); however, the TL effect was substantially greater for consonant nonwords than for vowel nonwords with the slowest responses (.9 quantile: 213 vs. 86 ms, respectively; interaction: $F(1, 24) = 6.18$, $p.025$).

⁴ This experiment was replicated with a shorter stimulus duration (200 ms, followed by a pattern mask composed of a series of # signs). The pattern of results was essentially the same, except that error rates were slightly higher. For the transposed consonant conditions, error rates were 46.3 vs. 9.0% for the transposed-letter nonwords and their orthographic controls (the mean RTs were 865 vs. 779 ms, respectively). For the transposed vowel conditions, error rates were 30.5 vs. 9.0% for the transposed-letter nonwords and their orthographic controls (the mean RTs were 846 vs. 773 ms, respectively). That is, the transposed-letter similarity effects were noticeably greater for the transposition of consonants than for the transposition of vowels, especially in the error data (37.3 vs. 21.5%; 86 vs. 73 ms).

General discussion

The present experiments allow the following conclusions: (1) nonword primes created by transposing two nonadjacent letters produce masked priming effects relative to the appropriate orthographic controls (*caniso-CASINO* vs. *caviro-CASINO*), (2) these priming effects seem to be restricted to the case in which the transposed letters are consonants (i.e., *anamil-ANIMAL* is no faster than the orthographic control *anomal-ANIMAL*), and (3) in a single-presentation lexical decision task, both TL-consonant nonwords (*TRADEGIA*) and TL-vowel nonwords (*ABSULOTO*) produce longer latencies and more errors than control nonwords (*TRABEPIA*, *ABSELITO*), however, TL-consonant nonwords are somewhat more problematic (as indicated by the false positive error rates) than TL-vowel nonwords.

The presence of nonadjacent TL similarity effects poses a clear problem for models that assume a position-specific coding scheme (e.g., the interactive-activation model and its extensions). As stated in the Introduction, the only real way to try to explain an effect of this sort is to incorporate the notion of noise in the coding process, so that the representation of one letter is not immediately tied to a single letter position but, instead, extends activation into nearby letter positions. That is, the letter *S* in the nonadjacent TL nonword *CANISO* would provide considerable activation to the representation of the letter *S* in letter position 5, somewhat less activation to the representations of the letter *S* in adjacent positions (4 and 6), as well as at least some activation to the representation of the letter *S* in letter position 3.

Ratcliff's (1981) letter coding model is a model based on these ideas. In Ratcliff's model, each letter in a letter string creates a distribution of activation over positions. As a result, the representation of a letter in a given position would be activated by the appearance of that letter in any nearby letter position. Therefore, the overall overlap between *CASINO* and the TL nonwords *CANISO* would be substantially higher than that between *CASINO* and the orthographic control *CAVIRO* (see Gómez, Perea, & Ratcliff, 2003).

Even if Ratcliff's scheme were integrated into the interactive-activation architecture, however, new problems may emerge. That is, as Davis (1999) has noted, noise of this sort in the visual input code would harm the ability of IA-based models to recognize highly familiar inputs. Thus, it would be extremely unlikely that IA models with this modified coding scheme would still be able to successfully simulate many of the effects they now can. Further, there would, of course, be an additional problem with respect to the present data. An account of this sort has no way to distinguish between vowel and consonant processing.

An alternative way of explaining nonadjacent TL priming would be to assume either a spatial coding

scheme (as in the SOLAR model) or a letter-tagging coding scheme (as in the SERIOL model). Both of these coding schemes can readily capture the overall pattern of TL effects in the present experiments. That is, in both cases, the similarity between the nonadjacent TL nonwords and their corresponding base words is higher than the similarity between the orthographic controls (the two-letter different nonwords) and their corresponding base words (see Introduction), leading to the prediction of nonadjacent TL priming effects. (Similar predictions can be made by very recently proposed letter-coding models: the open-bigram model, Grainger & van Heuven, 2003; and the overlap model, Gómez et al., 2003.) The problem, however, is still that, like a modified position-specific model, the current versions of all these models do not make any distinctions between the processing of vowels and consonants. Thus, some of the model assumptions would need to be altered to accommodate the presence of priming effects from nonadjacent TL nonwords based on consonants but not on vowels (Experiments 1–3) and the presence of a substantially larger number of false positive errors in a single-presentation lexical decision task for TL-consonant nonwords than for TL-vowel nonwords (Experiment 4).

To allow these models to explain the present data, one would need to start by assuming that very early in processing, the consonant/vowel status of a given letter is obtained and that this information allows the two types of letters to be segregated from each other. For example, some formulations (e.g., Berent et al., 2001; Caramazza & Miceli, 1990; Tainturier & Caramazza, 1996) assume that very early in processing, the word *CASINO* activates a C.V.C.V.C.V structure with the identified letters then filling those slots. This segregation of consonants from vowels, however it is accomplished, would then allow subsequent processing of the consonant letters (*c · s · n*) to differ from the processing of the vowel letters (*a · i · o*) (e.g., see Caramazza et al. (2000) and Nespor et al. (2003) for discussions of how consonants and vowels might play different roles in lexical processing). One would next need to assume that the different processing that consonants and vowels underwent made it somewhat more likely that one would observe writing/reading errors involving transpositions of two consonants than transpositions of two vowels. Note that one implication of these assumptions would be that it should be very unusual to find errors involving transpositions of one consonant and one vowel. Indeed, the usual pattern obtained with brain damaged patients is the transposition/replacement of vowels by vowels, and consonants by consonants (see Caramazza et al., 2000).

One way in which a segregation of consonants and vowels could occur would be for the orthographic coding of vowels to be more distinct than that of consonants. For example, in Ratcliff's (1981; Gómez et al.,

2003) framework, the activation gradients corresponding to vowels and consonants could be assumed to differ (see Drevnowski, 1980, for a similar suggestion). More specifically, Ratcliff's model could explain the present findings by assuming that the activation of consonant graphemes extends further (i.e., into nearby letter slots) than the activation of vowel graphemes.⁵

Within the framework of the SERIOL or SOLAR models, the key aspect of the models would be the set of activation levels of the sublexical units. As noted, the SOLAR model uses activation levels to code order information (i.e., the first letter is coded by the highest activation value, the second letter is coded with a slightly smaller activation value, etc.). According to the model, there is a constant ratio (e.g., 2:1) between the activation levels of successive letters (the so-called the *invariance principle*, see Davis, 1999), without any distinctions between consonants and vowels. The SERIOL model works in a fairly similar way. In that model (Whitney, 2001; Whitney & Berndt, 1999), once letter positions have been tagged, the activation of a given letter node is set to $0.7^{\text{position}-1}$. Thus, for the word *CASINO*, the activation levels for its component letters would be C-1, A-0.7, S-0.49, I-0.34, N-0.24, and O-0.17, with, again, no distinction between consonants and vowels.

If, as suggested above, consonant/vowel status is ascertained very early in processing, these models could propose that the letters achieve different activation levels depending on whether they are consonants or vowels. Probably, the most reasonable way to do this would be either to propose a heightened activation level if the letter were a vowel (making it more distinct) or a lowered activation level if a letter were a consonant (making it less distinct). For example, in SERIOL's framework, the consonant *S* in *CASINO* could have a reduced activation level of 0.40, which would make its level more similar to that of the *N* (although the activation level of the *N* would, presumably, also be lowered). Thus, the *S* and the *N* would be more likely to show TL effects than the *A* and the *I* (i.e., *CANISO* would be more similar to *CASINO* than *CISANO* would). A similar fix for SOLAR would lead to similar predictions.

There are, of course, limits to the degree to which activation levels can be changed without altering the basic structure of the models. For example, the activation level for the consonant *S* in position 3 in *CASINO* cannot be reduced below the activation level for the vowel *I* in position 4 or else the activation levels would

no longer accurately code letter order. Whether changes to the models within these limits would actually allow the models to predict the vowel/consonant differences and whether these changes would then harm the models' abilities to explain other letter coding results would, of course, be questions for future research.

An additional point that should be noted is that SERIOL, unlike SOLAR, incorporates a bigram level. Thus, in theory, SERIOL has the potential to explain the vowel/consonant differences in terms of activation patterns at that level, rather than at the letter level. For example, vowel bigrams (i.e., the *AI* pair, the *AO* pair, and the *IO* pair in *CASINO*) could be assumed to have higher activation levels than consonant bigrams (e.g., *SN*). In this way, *CISANO* would not be very similar to *CASINO* because the set of bigram codes activated by *CISANO* would not include the crucial *AI* pair. Hence, *CISANO* would not be expected to produce much priming of *CASINO*.

Alternatively, it would be possible to propose a locus of the consonant/vowel differences observed here which is entirely outside of the architecture of the models under consideration. For example, one could propose that those differences arise at the sub-lexical *phonological* level.⁶ Subjectively, the transposition of two consonants does appear to preserve more of the sound of the original word than the transposition of two vowels (e.g., compare the TL-consonant nonword *LIREBACIÓN* to its base word, *LIBERACIÓN*, in contrast to the TL-vowel nonword *LIBARECIÓN*). Indeed, vowel sounds seem to become phonologically relevant earlier in life than consonant sounds (see Bertoncini, Bijeljac-Babic, Jusczyk, Kennedy, & Mehler, 1988) and (at least in Spanish) young children spell the vowel sounds before consonant sounds (e.g., *A_I_O_A* for *MARIPOSA*, the Spanish for butterfly; see Ferreiro & Teberosky, 1982). If the masked priming effects we have observed are phonological effects (see Hino, Lupker, Ogawa, & Sears, 2003), the obvious prediction would be that TL-consonant primes would be better primes than TL-vowel primes because TL-vowel primes would be more dissimilar to their targets than TL-consonant primes. In contrast, if negative lexical decision responses are based more on an analysis of the orthographic structure of the letter string, it would make sense that the consonant/vowel difference might diminish. Further research and simulation work is, of course, needed to examine this possibility.

With respect to the general issue of vowel-consonant differences and the impact of phonological codes, it is worth noting that Berent and Perfetti (1995) have proposed that because consonants and vowels are distinct linguistic entities (with the argument being that this is a universal structural distinction; see also Caramazza

⁵ A possible reason why the activation gradient for consonants might be broader than the activation gradient for vowels could be simply because consonants are less frequent than vowels (rather than being a function of basic structural differences). Although this is an unlikely explanation (e.g., see Caramazza et al., 2000; for discussion), it is, potentially, a testable hypothesis.

⁶ We thank Carol Whitney for this suggestion.

et al., 2000), the time course of phonological assembly for vowels and consonants differs. In particular, Berent and Perfetti have produced evidence that consonants are processed faster than vowels in English (see also Lee et al., 2001).

As just discussed, the present findings do support the claim that there is an essential structural distinction between vowels and consonants. However, an explanation based on Berent and Perfetti's specific claim about the time course of phonological assembly for vowels and consonants is not likely to be relevant here. Instead, it seems more likely that the time course of vowel versus consonant processing would actually vary as a function of the characteristics of a given language. Specifically, in English, the grapheme–phoneme relationship for consonants is much more consistent than it is for vowels (Brown & Besner, 1987; Carr & Pollatsek, 1985). Thus, it would follow that English consonants would be coded more rapidly than English vowels. In contrast, in Spanish and in other Romance languages the grapheme–phoneme relationship for vowels is at least as consistent as that for consonants. Thus, vowel coding should occur at least as rapidly for vowels as for consonants. It is not at all surprising, therefore, that Colombo, Zorzi, Cubelli, and Brivio (2003) were unable to replicate Berent and Perfetti's consonant advantage using Italian.

Note also that, in spite of the fact that English and Spanish differ dramatically in terms of the consistency of phonological coding of vowels, the present results do have parallels in English. As noted, Perea and Lupker (2003b), using English stimuli, also found substantially larger TL priming effects when the transposition involved consonants rather than vowels. In addition, an experiment parallel to Experiment 4 carried out with English stimuli in the second author's lab yielded results that completely parallel Experiment 4's results (i.e., TL similarity effects involving nonadjacent letter positions are greater with consonants than with vowels when considering the nonwords in a lexical decision task). It seems likely, therefore, that the basis of the consonant/vowel differences is something that is common to the two languages. Thus, what we have attempted to do is to try to explain these differences in terms of a general principle (e.g., establishing the exact position of vowels may be more important for a word's identification than establishing the exact position of consonants) that would be relevant in both languages.

To summarize, we have provided empirical evidence that TL similarity effects occur even when the transposed letters are not adjacent. These results are as predicted by models with a spatial or letter-tagging coding scheme (SOLAR and SERIOL models). In addition, however, we discovered that this conclusion applies especially to consonants, supporting the claims that there are some basic processing differences between vowels and consonants (see Caramazza et al., 2000). This latter

fact cannot be readily explained by either the SOLAR or SERIOL models. Whether those models can be amended to allow them to explain the present results, as well as consonant/vowel differences in general, is a question for future research.

Appendix. Related pairs in Experiments 1–3

Experiment 1

The items are arranged in quadruplets in the following order: one-letter different prime, TL-prime, RL-prime, and target word.

camafa, cadama, cafasa, CAMADA; gaceda, gateca, gadena, GACETA; sonifo, sodino, sofivo, SONIDO; morabo, modaro, mobaso, MORADO; felivo, fenilo, fevito, FELINO; dibupo, dijubo, diputo, DIBUJO; humaro, hunamo, huraco, HUMANO; romaso, ronamo, rosavo, ROMANO; manesa, marena, maseva, MANERA; masape, majase, mapane, MASAJE; basuva, barusa, bavuna, BASURA; malefa, matela, mafeda, MALETA; meduva, mesuda, mecuta, MEDUSA; cadesa, caneda, casela, CADENA; visila, vitisa, vilica, VISITA; genosa, gemona, gesova, GENOMA; tiravo, tinaro, tivaso, TIRANO; fabata, fadaba, fatala, FABADA; gemeto, gelemo, geteco, GEMELO; figuca, firuga, ficupa, FIGURA; mirata, midara, mitasa, MIRADA; parato, padaro, pataso, PARADO; dorato, dodaro, dotaso, DORADO; tutefa, tuleta, tufeba, TUTELA; lativo, lanito, lavido, LATINO; galoje, gapole, gajode, GALOPE; regafo, relago, refapo, REGALO; navaya, najava, nayasa, NAVAJA; corosa, conora, cosova, CORONA; veciso, venico, vesiso, VECINO; debapo, dejabo, depafo, DEBAJO; modebo, moledo, mobeto, MODELO; monela, modena, molesa, MONEDA; minudo, mituno, miduso, MINUTO; pecafo, pedaco, pefano, PECADO; golono, gosolo, gonoto, GOLOSO; pesato, pedaso, petano, PESADO; gusazo, gunaso, guzavo, GUSANO; rulefa, rutela, rufeda, RULETA; dineso, direno, diseco, DINERO; diviro, dinivo, dirico, DIVINO; camiva, casima, cavica, CAMISA; rodape, rojade, ropate, RODAJE; camivo, canimo, cavico, CAMINO; divima, disiva, dimica, DIVISA; filede, fitele, fidebe, FILETE; relado, retalo, redafo, RELATO; pasino, pavisio, panino, PASIVO; butava, bucata, buvala, BUTACA; sumivo, susimo, suvino, SUMISO; cadeva, careda, cavela, CADERA; cubaso, cunabo, cusato, CUBANO; decaso, denaco, desavo, DECANO; tariba, tafira, tabima, TARIFA; madeva, mareda, maveta, MADERA; senafo, sedano, sefavo, SENADO; noveta, noleva, noteca, NOVELA; rutiva, runita, ruvifa, RUTINA; ganato, gadano, gatato, GANADO; nacito, nadico, natiso, NACIDO; pasato, padaso, patano, PASADO; comifa, codima, cofisa, COMIDA; famono, fasomo, fanovo, FAMOSO; gotena, goreta, goned, GOTERA; lujomo, lusojo, lumopo, LUJOSO; tomaf, totame, tofave, TOMATE; cocifo, codico, cofino, COCIDO; zarifo, zanito, ZAFIRO; docesa, doneca, dosema, DOCENA; tocabo, todaco, tobaso, TOCADO; marifo, madiro, mafino, MARIDO; facela, fateca, falena, FACETA; polaso, pocalo, posato, POLACO; racino, ramico, ranino, RACIMO; herifa, hedira, hefiva, HERIDA; jugomo, jusogo, jumopo, JUGOSO; harisa, hanira, hasima, HARINA; herepe, hejere, hepese, HEREJE; salubo, sadulo, sabufo, SALUDO; canefa, calena, cafesa, CANELA; moliso, monilo, mosifo, MOLINO; futuvo,

furuto, fuvubo, FUTURO; sonoso, sorono, sosoco, SONORO; ligeso, lirego, liseyo, LIGERO; tesono, teroso, tenovo, TESORO; boleno, borelo, boneto, BOLERO; sucemo, suseco, sumevo, SUCESO; marivo, maniro, mavico, MARINO; veraco, venaro, vecaso, VERANO; rotuna, roruta, ronula, ROTURA; rebapa, rejaba, repata, REBAJA; seguso, serugo, sesuyo, SEGURO; tabavo, tacabo, tavato, TABACO; pedano, pezado, penato, PEDAZO; sonefo, soteno, sofeco, SONETO; gomisa, gonima, gosisa, GOMINA; helafo, hedalo, hefabo, HELADO; salifa, sadila, safita, SALIDA; coyobe, cotoye, coboge, COYOTE; semasa, senama, sesava, SEMANA; casiro, caniso, caviro, CASINO; butaso, bunato, busafo, BUTANO; moromo, mosoro, momovo, MOROSO; malago, mayalo, magato, MALAYO; bonifo, botino, bofivo, BONITO; sujedo, sutejo, sudeyo, SUJETO; gorita, golira, gotina, GORILA; peloda, petola, pedofa, PELOTA; corava, cozara, covasa, CORAZA; celoro, cesolo, cerofo, CELOSO; bigole, bitoge, bilope, BIGOTE; vasipa, vajisa, vapira, VASIJA; balata, badala, batafa, BALADA; motino, movito, monido, MOTIVO; lejavo, lenajo, levapo, LEJANO; menuto, meduno, metuco, MENUDO; ranuva, raruna, ravusa, RANURA; tapele, tatepe, talege, TAPETE; casato, cadaso, catano, CASADO; mesefa, metesa, mefeve, MESETA; bikiri, biniki, bisili, BIKINI; payamo, pasayo, pamago, PAYASO; cabena, cazeza, caneta, CABEZA; coliva, conila, coviba, COLINA; rabivo, ranibo, ravito, RABINO; repima, resipa, remiga, REPISA; parape, pajare, papase, PARAJE; cobaga, coyaba, cogala, COBAYA.

Experiment 2

The items are arranged in quadruplets in the following order: one-letter different prime, TL-prime, RL-prime, and target word.

agujes, agajus, agejos, AGUJAS; alemón, alámen, alímun, ALEMÁN; animol, anamil, anomal, ANIMAL; ruides, ruodis, ruades, RUIDOS; piedod, piaded, piudod, PIEDAD; apoyur, apayer, APOYAR; huesas, huoses, huasis, HUESOS; elegar, eliger, elogar, ELEGIR; chalut, chelat, chilot, CHALLET; avidaz, avediz, avoduz, AVIDEZ; glorua, gliroa, glerua, GLORIA; ademís, adámes, adímos, ADEMÁS; agonúa, agínoa, agénua, AGONÍA; vuelis, vuoles, vuilas, VUELOS; prevoa, privea, pruvoa, PREVIA; emisar, emosir, emusar, EMISOR; frágel, frigál, frogél, FRÁGIL; náusoa, náesua, náisua, NÁUSEA; azotúa, azetoe, azitua, AZOTEA; éxitas, éxotis, éxutas, ÉXITOS; planis, plonas, plines, PLANOS; acudar, acidur, acedor, ACUDIR; llegor, llager, llogir, LLEGAR; cuevos, cuaves, cuivos, CUEVAS; tregoa, trugea, trigoa, TREGUA; chófar, chefór, chifar, CHÓFER; llenor, llaner, llunor, LLENAR; árabos, árebas, árubos, ÁRABES; plurel, plarul, plerol, PLURAL; acidaz, acediz, acudoz, ACIDEZ; plumis, plamus, plumas, PLUMAS; nietus, niotes, niatus, NIETOS; globel, glabol, glibel, GLOBAL; criman, cremin, cramon, CRIMEN; amiges, amogis, amuges, AMIGOS; asumer, asimur, asemor, ASUMIR; inútel, initúl, inatél, INÚTIL; viejis, viejos, viujas, VIEJOS; unided, unadid, unedod, UNIDAD; travós, trévas, trívos, TRAVÉS; lloer, llorar, llurer, LLORAR; ademin, adámen, adómin, ADEMÁN; florus, fleros, fliras, FLORES; placir, plecra, plicor, PLACER; evitor, evatir, evutor, EVITAR; ciuded, ciadud, ciedod, CIUDAD; premuo, primeo, prumao, PREMIO; cruzor, crazur, crezir, CRUZAR;

ácidus, ácodis, ácedus, ÁCIDOS; chicus, chacis, choces, CHICAS; cuider, cuadir, cuoder, CUIDAR; origon, oregon, orugan, ORIGEN; ciegis, cioges, ciagus, CIEGOS; lluvea, llivua, llevoa, LLUVIA; examon, exeman, eximon, EXAMEN; ayudor, aya-dur, ayedor, AYUDAR; añador, añidar, añedor, AÑADIR; reunor, reinur, reanor, REUNIR; clasos, clesas, clisos, CLASES; brutol, bratul, bretil, BRUTAL; imitor, imatir, imetor, IMITAR; clamir, clomar, clumer, CLAMOR; azúcor, azacúr, azecír, AZÚCAR; cránuo, crenáo, crinúo, CRÁNEO; ayuner, ayanur, ayenor, AYUNAR; deudar, deodur, deadir, DEUDOR; chinus, chanis, chenos, CHINAS; cremus, crames, cri-mos, CREMAS; orader, orodar, oruder, ORADOR; guñal, guoñil, guañel, GUÑOL; suizes, suozis, suezas, SUIZOS; ahogur, ahagor, ahiger, AHOGAR; grutos, gratus, grotes, GRUTAS; ilegol, ilagel, ilugol, ILEGAL; clavol, cleval, clivol, CLAVEL; oculor, ocalur, ocelir, OCULAR; asesir, asoser, asusir, ASESOR; afiner, afanir, afenor, AFINAR; anemoa, anínea, anumoa, ANEMIA; acosur, acasor, acuser, ACOSAR; abulea, abilua, abeloa, ABULIA; bribén, bróbin, bríban, BRIBÓN; agitor, agatir, agetor, AGITAR; grupol, grapul, gripel, GRUPAL; apuror, aparur, apirer, APURAR; adaguo, adigao, adugeo, ADAGIO; asoler, asalor, aselir, ASOLAR; aflor, afalir, afelor, AFILAR; clonir, clanor, cliner, CLONAR; florel, flarol, fleril, FLORAL; erizus, erozis, eruzas, ERIZOS; grasus, grosas, gruses, GRASOS; acunor, acañur, aciñer, ACUÑAR; fritus, frotis, frates, FRITOS; tribel, trabil, trubel, TRIBAL; brujis, brojus, brajes, BRUJOS; gruñor, griñur, greñar, GRUNIR; asader, asodar, asidur, ASADOR; aludor, alidur, aledar, ALUDIR; crátor, cretár, crítór, CRÁTER; acogur, acegor, acigur, ACOGER; aboler, abilor, abelur, ABOLIR; troter, trator, triter, TROTAR; dietos, diates, diutos, DIETAS; acoter, acator, acuter, ACOTAR; peatín, peótan, peítón, PEATÓN; apagín, apógan, apúgen, APAGÓN; olivus, olovís, olevas, OLIVOS; llanus, llonas, llines, LLANOS; alegor, alager, aligor, ALEGAR; abetus, abotes, abutis, ABETOS; trufis, trafus, trifes, TRUFAS; averúa, avírea, avúroa, AVERÍA; reinor, reanir, reunor, REINAR; grífes, grofís, gre-fas, GRIFOS; peínor, peanir, peonur, PEINAR; flacus, flocas, flices, FLACOS; evadur, evider, evudor, EVADIR; acusor, acasur, acosir, ACUSAR; anotir, anator, aniter, ANOTAR; chozus, chazos, chizes, CHOZAS; crudis, crodus, crades, CRUDOS; charel, choral, cherul, CHAROL; suécis, suoces, suacis, SUECOS; apelor, apaler, apolir, APELAR; frutol, fratul, fritel, FRUTAL; alojer, alajor, alijer, ALOJAR; orugos, oragus, origes, ORUGAS.

Experiment 3

The items are arranged in triplets in the following order: TL-prime, RL-prime, and target word.

Consonant transpositions/replacements: revuloción, revaleción, REVOLUCIÓN; evedinte, evadonte, EVIDENTE; enomorado, enimurado, ENAMORADO; impisoble, impus-able, IMPOSIBLE; inivetable, inovatable, INEVITABLE; amenacer, amonicer, AMANECER; genareción, genurición, GENERACIÓN; difuciltad, difoceltad, DIFICULTAD; escele-lara, escilora, ESCALERA; nataruleza, nateroleza, NATU-RALEZA; horozinte, horuzente, HORIZONTE; adalente, adolinte, ADELANTE; anamiles, anomeles, ANIMALES; libareción, liburición, LIBERACIÓN; femineño, femonuno,

FEMENINO; comasirio, comuserio, COMISARIO; semajente, semijonte, SEMEJANTE; avineda, avonuda, AVENIDA; opareción, oporución, OPERACIÓN; ilisuón, ilosaón, ILUSIÓN; habatición, haboteción, HABITACIÓN; unofirme, unaferme, UNIFORME; telivesión, teluvación, TELEVISIÓN; amirallo, amorello, AMARILLO; cominidad, comenodad, COMUNIDAD; aperace, aporuce, APARECE; litaretura, litorotura, LITERATURA; amaneza, amonuza, AMENAZA; amirecanos, amoracanos, AMERICANOS; enimego, enumago, ENEMIGO; autirodad, auturadad, AUTORIDAD; velicodad, velucadad, VELOCIDAD; unevirso, unovarsó, UNIVERSO; únacimente, únecomente, ÚNICAMENTE; absuloto, absileto, ABSOLUTO; señirota, señurata, SEÑORITA; imiganar, imogenar, IMAGINAR; expisoción, expasución, EXPOSICIÓN; evuloción, evileción, EVOLUCIÓN; opisoción, opusación, OPOSICIÓN.

Vowel transpositions/replacements: silimar, sitinar, SIMILAR; gusbata, gusdala, GUSTABA; namiciente, navipiento, NACIMIENTO; crícita, crírita, CRÍTICA; condiserar, conticerrar, CONSIDERAR; dormitorio, dorlinorio, DORMITORIO; esdutio, esbulio, ESTUDIO; doroles, dovotes, DOLORES; cazebras, caredas, CABEZAS; coroles, covotes, COLORES; sotilario, sofibarrio, SOLITARIO; hatibual, halidual, HABITUAL; hernamo, hersaso, HERMANO; prosópito, pronógito, PROPÓSITO; anretior, ancelior, ANTERIOR; esmótago, esnófago, ESTÓMAGO; descapio, desragio, DESPACIO; platena, plafema, PLANETA; rebúplica, redúglica, REPÚBLICA; tradegia, trabepia, TRAGEDIA; anásilis, anáritis, ANÁLISIS; posatio, pozalio, POTASIO; sonsira, soncina, SONRISA; tracidióñ, trasibióñ, TRADICIÓN; zatapos, zabagos, ZAPATOS; nodevad, nobesad, NOVEDAD; ténnimos, térvisos, TÉRMINOS; canimos, carisos, CAMINOS; prosefor, procetor, PROFESOR; ornedador, orcebador, ORDENADOR; juscitia, jusrilia, JUSTICIA; esrepanza, esnebanza, ESPERANZA; deyasuno, dejavuno, DESAYUNO; coditiana, cobiliana, COTIDIANA; dispiclina, disgiblina, DISCIPLINA; oldivado, oltsiado, OLVIDADO; desduno, desburo, DESNUDO; cladirad, clatinad, CLARIDAD; convidiód, concibiód, CONDICIÓN; sencivio, sernisio, SERVICIO.

Note that the nonword primes in Experiment 3 were the nonword targets in Experiment 4 (single-presentation lexical decision task).

References

- Alameda, J. R., & Cuetos, F. (1995). *Diccionario de frecuencia de las unidades lingüísticas del castellano*. Oviedo: Universidad de Oviedo.
- Andrews, S. (1996). Lexical retrieval and selection processes: Effects of transposed-letter confusability. *Journal of Memory and Language*, 35, 775–800.
- Berent, I., & Perfetti, C. A. (1995). A rose is a REEZ: The two cycles model of phonology assembly in reading English. *Psychological Review*, 102, 146–184.
- Berent, I., Bouissa, R., & Tuller, B. (2001). The effect of shared structure and content on reading nonwords: Evidence for a CV skeleton. *Journal of Experimental Psychology: Learning, Memory, Cognition*, 27, 1042–1057.
- Bertoncini, J., Bijeljac-Babic, R., Jusczyk, P. W., Kennedy, L. J., & Mehler, J. (1988). An investigation of young infants' perceptual representation of speech sounds. *Journal of Experimental Psychology: General*, 117, 21–33.
- Boatman, D., Hall, C., Goldstein, M. H., Lesser, R., & Gordon, B. (1997). Neuropceptual differences in consonant and vowel discrimination: As revealed by direct cortical electrical interference. *Cortex*, 33, 83–98.
- Brown, P., & Besner, D. (1987). The assembly of phonology in oral reading: A new model. In M. Coltheart (Ed.), *Attention and performance XII: The psychology of reading* (pp. 471–489). Hillsdale, NJ: Erlbaum.
- Caramazza, A., Chialant, D., Capasso, D., & Miceli, G. (2000). Separable processing of consonants and vowels. *Nature*, 403, 428–430.
- Caramazza, A., & Miceli, G. (1990). The structure of graphemic representations. *Cognition*, 37, 243–297.
- Carr, T. H., & Pollatsek, A. (1985). Recognizing printed words: A look at current models. In D. Besner, T. Waller, & G. E. Mackinnon (Eds.), *Reading research: Advances in theory and practice* (Vol. 5, pp. 1–82). Orlando, FL: Academic Press.
- Carreiras, M., & Perea, M. (2004). Naming pseudowords in Spanish: Effects of syllable frequency. *Brain and Language*, 90, 393–400.
- Colombo, L., Zorzi, M., Cubelli, R., & Brivio, C. (2003). The status of consonants and vowels in phonological assembly: Testing the two-cycles model with Italian. *European Journal of Cognitive Psychology*, 15, 405–433.
- Coltheart, M., Davelaar, E., Jonasson, J. F., & Besner, D. (1977). Access to the internal lexicon. In S. Dornic (Ed.), *Attention and performance VI* (pp. 535–555). Hillsdale, NJ: Erlbaum.
- Coltheart, M., Rastle, K., Perry, C., Ziegler, J., & Langdon, R. (2001). DRC: A dual-route cascaded model of visual word recognition and reading aloud. *Psychological Review*, 108, 204–256.
- Chambers, S. M. (1979). Letter and order information in lexical access. *Journal of Verbal Learning and Verbal Behavior*, 18, 225–241.
- Cubelli, R. (1991). A selective deficit for writing vowels in acquired dysgraphia. *Nature*, 353, 258–260.
- Davis, C. J. (1999). *The self-organising lexical acquisition and recognition (SOLAR) model of visual word recognition*. Unpublished doctoral dissertation, University of New South Wales.
- Drewnowski, A. (1980). Memory functions for vowels and consonants: A reinterpretation of acoustic similarity effects. *Journal of Verbal Learning and Verbal Behavior*, 19, 176–193.
- Ferreiro, E., & Teberosky, A. (1982). *Literacy before schooling*. Portsmouth, NH: Heinemann.
- Ferreres, A. R., López, C. V., Petracci, B., & China, N. N. (2000). Alexia por alteración de la vía perilexical de lectura. *Revista Neurológica Argentina*, 25, 17–28.
- Forster, K. I., & Davis, C. (1984). Repetition priming and frequency attenuation in lexical access. *Journal of Experimental Psychology: Learning, Memory, Cognition*, 10, 680–698.
- Forster, K. I., Davis, C., Schoknecht, C., & Carter, R. (1987). Masked priming with graphemically related forms:

- Repetition or partial activation? *Quarterly Journal of Experimental Psychology*, 39, 211–251.
- Forster, K. I., & Forster, J. C. (2003). DMDX: A Windows display program with millisecond accuracy. *Behavior Research Methods, Instruments, and Computers*, 35, 116–124.
- Forster, K. I., & Shen, D. (1996). No enemies in the neighborhood: Absence of inhibitory neighborhood effects in lexical decision and semantic categorization. *Journal of Experimental Psychology: Learning, Memory and Cognition*, 22, 696–713.
- Gafos, D. (1998). Eliminating long-distance consonantal spreading. *Natural Language and Linguistic Theory*, 16, 223–278.
- Goldsmith, J. (1990). *Autosegmental and metrical phonology*. Cambridge, UK: Basil Blackwell.
- Gómez, P., Perea, M., & Ratcliff, R. (2003, November). The overlap model of the encoding of letter positions. Poster presented at the 45th Annual Meeting of the Psychonomic Society, Vancouver, BC, Canada.
- Grainger, J., & Jacobs, A. M. (1996). Orthographic processing in visual word recognition: A multiple read-out model. *Psychological Review*, 103, 518–565.
- Grainger, J., & van Heuven, W. J. B. (2003). Modeling letter position coding in printed word perception. In P. Bonin (Ed.), *The mental lexicon* (pp. 1–23). New York: Nova Science.
- Hino, Y., Lupker, S. J., Ogawa, T., & Sears, C. R. (2003). Masked repetition priming and word frequency effects across different types of Japanese scripts: An examination of the lexical activation account. *Journal of Memory and Language*, 48, 33–66.
- Kay, J., & Hanley, R. (1994). Peripheral disorders of spelling: The role of the graphemic buffer. In G. D. A. Brown & N. C. Ellis (Eds.), *Handbook of spelling: Theory, process and intervention* (pp. 295–315). Chichester: Wiley.
- Lane, D. M., & Ashby, B. (1987). PsychLib: A library of machine language routines for controlling psychology experiments on the Apple Macintosh computer. *Behavior Research Methods, Instruments, and Computers*, 19, 246–248.
- Lee, H.-W., Rayner, K., & Pollatsek, A. (2001). The relative contribution of consonants and vowels to word identification during reading. *Journal of Memory and Language*, 44, 189–205.
- Lee, H.-W., Rayner, K., & Pollatsek, A. (2002). The processing of consonants and vowels in reading: Evidence from the fast priming paradigm. *Psychonomic Bulletin & Review*, 9, 766–772.
- Mewhort, D. J. K., Campbell, A. J., Marchetti, F. M., & Campbell, J. I. D. (1981). Identification, localization, and “iconic memory”: An evaluation of the bar-probe task. *Memory & Cognition*, 9, 50–67.
- Monaghan, P., & Shillcock, R. (2003). Connectionist modeling of the separable processing of consonants and vowels. *Brain and Language*, 86, 83–98.
- Mozer, M. (1987). Early parallel processing in reading: A connectionist approach. In M. Coltheart (Ed.), *Attention and performance XII: The psychology of reading* (pp. 83–104). Hillsdale, NJ: Erlbaum.
- Nespor, M., Peña, M., & Mehler, J. (2003). On the different roles of vowels and consonants in speech processing and language acquisition. *Lingue e Linguaggio*, 2, 221–247.
- O'Connor, R. E., & Forster, K. I. (1981). Criterion bias and search sequence bias in word recognition. *Memory & Cognition*, 9, 78–92.
- Perea, M., & Carreiras, M. (1998). Effects of syllable frequency and syllable neighborhood frequency in visual word recognition. *Journal of Experimental Psychology: Human Perception and Performance*, 24, 134–144.
- Perea, M., & Lupker, S. J. (2003a). Does jugde activate COURT? Transposed-letter similarity effects in masked associative priming. *Memory & Cognition*, 31, 829–841.
- Perea, M., & Lupker, S. J. (2003b). Transposed-letter confusability effects in masked form priming. In S. Kinoshita & S. J. Lupker (Eds.), *Masked priming: State of the art* (pp. 97–120). Hove, UK: Psychology Press.
- Perea, M., & Rosa, E. (2000). Repetition and form priming interact with neighborhood density at a brief stimulus-onset asynchrony. *Psychonomic Bulletin & Review*, 7, 668–677.
- Peressotti, F., & Grainger, J. (1995). Letter-position coding in random consonant arrays. *Perception and Psychophysics*, 57, 875–890.
- Peressotti, F., & Grainger, J. (1999). The role of letter identity and letter position in orthographic priming. *Perception and Psychophysics*, 61, 691–703.
- Pérez, E., Palma, A., & Santiago, J. (2001). Una base de datos de errores del lenguaje en castellano [A database of speech errors in Spanish]. Poster presented at the V Simposio de Psicolingüística, Granada, Spain.
- Ratcliff, R. (1981). A theory of order relations in perceptual matching. *Psychological Review*, 88, 552–572.
- Rumelhart, D. E., & McClelland, J. L. (1982). An interactive activation model of context effects in letter perception: Part 2. The contextual enhancement effect and some tests and extensions of the model. *Psychological Review*, 89, 60–94.
- Schoonbaert, S., & Grainger, J. (2004). Letter position coding in printed word perception: Effects of repeated and transposed letters. *Language and Cognitive Processes*, 19, 333–367.
- Sears, C. R., Hino, Y., & Lupker, S. J. (1995). Neighborhood frequency and neighborhood size effects in word recognition. *Journal of Experimental Psychology: Human Perception and Performance*, 21, 876–900.
- Siakaluk, P. D., Sears, C. R., & Lupker, S. J. (2002). Orthographic neighborhood effects in lexical decision: The effects of nonword orthographic neighborhood size. *Journal of Experimental Psychology: Human Perception and Performance*, 28, 661–681.
- Tainturier, M. J., & Caramazza, A. (1996). The status of double letters in graphemic representations. *Journal of Memory and Language*, 35, 53–73.
- Westall, R., Perkey, M. N., & Chute, D. L. (1986). Accurate millisecond timing on the Apple Macintosh using Drexler's Millitimer. *Behavior Research Methods, Instruments, and Computers*, 18, 307–311.
- Whitney, C., & Berndt, S. L. (1999). A new model of letter string encoding: Simulating right neglect dyslexia. In J. A. Reggia, E. Rupp, & D. Glanzman (Eds.), *Progress in brain research* (Vol. 121, pp. 143–163). Amsterdam: Elsevier.
- 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.