

Does *jugde* activate COURT? Transposed-letter similarity effects in masked associative priming

MANUEL PEREA

Universitat de València, València, Spain

and

STEPHEN J. LUPKER

University of Western Ontario, London, Ontario, Canada

Transposed-letter (TL) nonwords (e.g., *jugde*) can be easily misperceived as words, a fact that is somewhat inconsistent with the letter-position-coding schemes employed by most current models of visual word recognition. To examine this issue further, we conducted four masked semantic/associative priming experiments, using a lexical decision task. In Experiment 1, the related primes could be words, TL-internal nonwords, or replacement-letter (RL) nonwords (e.g., *judge*, *jugde*, or *judpe*, respectively; the target would be COURT). Relative to an unrelated condition, masked TL-internal primes produced a significant semantic/associative priming effect, an effect that was only slightly smaller than the priming effect for word primes. No effect, however, was observed for RL-nonword primes. In Experiment 2, the TL-nonword primes were created by switching the two final letters of the primes (e.g., *judeg*). The results again showed a semantic/associative priming effect for word primes, but not for TL-final nonword primes or for RL-nonword primes. Experiment 3 replicated the associative/semantic priming effect for TL-internal nonword primes, with, again, no effect for TL-final nonword primes. Finally, Experiment 4 again failed to yield a priming effect for TL-final nonword primes. The implications of these results for the choice of a letter-position-coding scheme in visual word recognition models are discussed.

One issue that all models of visual word recognition in alphabetic orthographies must ultimately take a position on is *how* the human processing system encodes letter positions when creating internal orthographic representations. Furthermore, although the choice of a coding scheme might seem to be a secondary aspect of these models, it can have a large impact on a model's predictions (Andrews, 1996). For example, virtually all of the current models assume that the derived orthographic representation activates the lexical representations of formally similar words (see Coltheart, Rastle, Perry, Ziegler, & Langdon, 2001; Forster, 1976; Grainger & Jacobs, 1996; Johnson & Pugh, 1994; McClelland & Rumelhart, 1981; Norris, 1986; Paap, Newsome, McDonald, & Schvaneveldt, 1982; Seidenberg & McClelland, 1989). Which words should be

considered to be formally similar to the presented letter string, however, depends to a large extent on how the particular model codes letter positions.

In many current models of visual word recognition (e.g., the multiple read-out model, Grainger & Jacobs, 1996; the interactive-activation model, McClelland & Rumelhart, 1981; the activation-verification model, Paap et al., 1982),¹ letter position coding is assumed to be *channel specific*. That is, letters are assumed to be tagged to their positions in the orthographic representation of the presented word well before the identities of the letters have been encoded. The implication is that, within these types of models, the nonword *jugde* is no more similar to the word *judge* than is the nonword *junpe*, and is less similar to *judge* than is the nonword *judpe*, because *judpe* overlaps *judge* in four out of five letter positions, whereas *jugde* and *junpe* both overlap in only three out of five letter positions. In contrast, there are now a couple of models (SERIOL, Whitney, 2001; SOLAR, Davis, 1999) that use coding schemes in which *jugde* would be regarded as more similar to *judge* than is *judpe*. In these schemes, the existence of the *g* and the *d* in *jugde*, even though they are in the wrong letter positions, increases the similarity of *jugde* to *judge*, whereas the presence of an incorrect letter, the *p*, in *judpe* noticeably decreases its similarity to *judge*.

Nonwords such as *jugde*, referred to as *transposed-letter* (TL) nonwords, are the focus of the present investigation. In fact, there are now a number of studies in the literature suggesting that, at some level, TL nonwords are

This research was supported by Natural Sciences and Engineering Research Council (NSERC) of Canada Grant A6333 to S.J.L. and by Grant BSO2002-03286 from the Spanish Ministry of Science and Technology to M.P. Portions of the research described in this article were presented at the Masked Priming: State of the Art Conference, Sydney, Australia, April 2001. The authors thank Tamsen Taylor, Yolanda Yuen, Karen Hussey, and Vanessa Muzzin for their assistance in the testing of participants. We are grateful to Steve Joordens, two anonymous reviewers, and the action editor, Jay Rueckl, for very thorough criticism that improved the paper substantially. Correspondence concerning this article should be addressed to M. Perea, Departament de Metodologia, Facultat de Psicologia, Av. Blasco Ibáñez, 21, 46010-València, Spain (e-mail: mperea@uv.es) or to S. J. Lupker, Department of Psychology, University of Western Ontario, London, ON, N6A 5C2 Canada (e-mail: lupker@uwo.ca).

actually quite similar to their base word (at least as similar as nonwords such as *judpe*, which we refer to as *replacement-letter* [RL] nonwords; e.g., Andrews, 1996; Chambers, 1979; Holmes & Ng, 1993; O'Connor & Forster, 1981; Perea & Lupker, 2003; Taft & van Graan, 1998). The specific goal of the present research was to determine whether TL nonwords activate semantic/associative information from their base words, as indexed by semantic/associative priming effects in a masked-priming lexical decision task.

In the masked-priming technique, a forward-masked lowercase prime is presented briefly (for around 40–66 msec) and is subsequently replaced by the uppercase target (see Forster & Davis, 1984; Forster, Mohan, & Hector, 2003). A number of studies have shown that masked word primes activate semantic/associative information, as demonstrated by the presence of semantic/associative priming effects with this technique (i.e., responses to NURSE are faster when it is preceded by the prime *doctor* than when it is preceded by the prime *butter*; see Bodner & Masson, in press; Bourassa & Besner, 1998; de Groot & Nas, 1991; Gonnerman & Plaut, 2000; Perea & Gotor, 1997; Perea & Rosa, 2002a, 2002b; Sereno, 1991; Williams, 1996).

Our basic empirical question was whether TL nonword primes would also activate semantic/associative information, producing priming effects. At present, there is evidence that some types of briefly presented nonwords can activate semantic/associative information from their base word, which, in turn, facilitates responding to a related target word. For example, Lukatela and Turvey (1994) reported that pseudohomophone primes (e.g., *nale* would be the prime for FINGER) can produce a reliable associative priming effect. Lukatela, Carello, Savić, Urošević, and Turvey (1998) found that if Roman and Cyrillic characters were mixed in the same string of letters, masked nonword primes such as *robot* (the appropriate word prime would be *роѡот*) facilitated lexical decisions to associated targets (e.g., AUTOMAT; automaton, in English). (Note that the letter *r* does not exist in the Cyrillic alphabet, whereas the letter *ѡ* does not exist in the Roman alphabet.) Finally, and most relevant to the present investigation, Bourassa and Besner (1998) reported that masked RL nonwords (e.g., *oceln*–WAVES) can also activate semantic/associative information, producing priming effects.

With respect to TL nonwords used as primes, evidence indicates that they are effective as form primes. For example, Forster, Davis, Schoknecht, and Carter (1987, Experiment 1) found that, relative to the results obtained in a control condition with unrelated primes and targets, masked-priming effects were essentially equivalent for identity primes (*answer*–ANSWER, 64 msec) and for TL-nonword primes (*anwser*–ANSWER, 63 msec), whereas masked-priming effects were slightly smaller for RL-nonword primes (e.g., *antwer*–ANSWER, 52 msec). This result suggests that a TL prime may be as effective in activating the representation of the target word as the target word itself, at least in a masking context. Likewise, in a recent series of experiments, Perea and Lupker (2003) found

that TL primes (e.g., *ocaen*–OCEAN) also produced reliable priming effects when compared with orthographic control primes (e.g., *ocuon*–OCEAN). Thus, it does seem clear that TL nonwords can produce masked *form*-priming effects.

The question asked here, whether TL nonwords can also activate semantic/associative information from their base words, producing semantic/associative priming, goes one step further. That is, as has been argued by Masson and colleagues (Bodner & Masson, 1997; Masson & Isaak, 1999), form-priming effects may be sublexical effects (although see Forster et al., 1987; Frost, Forster, & Deutsch, 1997). That is, form-priming effects may be due to activation of the sublexical units used in the creation of the orthographic representation, rather than to activation of the lexical unit for the base word. If the form-priming effects produced by TL-nonword primes are, indeed, sublexical, the implication would be that those primes should not produce semantic/associative priming effects. On the other hand, if TL nonwords do produce semantic/associative priming effects, it would clearly indicate that TL nonwords were activating the lexical/semantic representations of their base words, reinforcing the models in which the orthographic representations produced by *jugde* and *judge* and, hence, the patterns of lexical activation those two letter strings produce, are quite similar (e.g., SERIOL model, Whitney, 2001; SOLAR model, Davis, 1999).

Transposing Internal Letters Versus External Letters

It is generally assumed that the quality of information about letters (both positional information and identity information) is better at the ends of the word than in the middle of the word (e.g., Estes, Allmeyer, & Reder, 1976; Friedmann & Gvion, 2001; Jordan, 1990; Perea, 1998), due to the fact that exterior letters suffer less lateral interference from neighboring letters. As a result, end letters have often been regarded as having a special status in word recognition models (see, e.g., Forster, 1976; Grainger, 1992; Humphreys, Evett, & Quinlan, 1990; Rumelhart & McClelland, 1982). Humphreys et al. (1990; see also Jacobs, Rey, Ziegler, & Grainger, 1998), for example, suggested that letters appear to be coded into an orthographic representation in which their positions as internal or external letters are marked. Consistent with this view, Perea and Lupker (2003) found that form-priming effects, relative to results with an orthographic control, were greater for TL-internal primes (*budget*–BUDGET relative to *buijet*–BUDGET) than for TL-final primes (*budgete*–BUDGET relative to *budgetfa*–BUDGET). Similarly, using a single-presentation lexical decision task, Chambers (1979) found that TL nonwords were more difficult to reject when they were constructed by switching two internal letters (e.g., *eviednce*) than when they were constructed by switching the two initial or two final letters (e.g., *amgazine* or *domestci*).

If the quality of letter position information is higher at the ends of the word than in the middle, nonwords such as *judge* may be less likely to activate the lexical representation for JUDGE than nonwords such as *jugde* would be, due

to the mismatch in the final letter position. If so, TL-final primes would be less likely to access semantic/associative information from their base words than TL-internal primes would be and, hence, less likely to produce semantic/associative priming effects. In order to examine these issues, the TL primes in Experiment 1 were created by transposing two internal letters of five-letter words (e.g., *jugde*), whereas the TL primes in Experiment 2 were created by transposing the two final letters of five-letter words (*judeg*). (Our decision to use TL-final primes, rather than TL-initial primes, was motivated mainly by the fact that we also, somewhat arbitrarily, chose to use TL-final, rather than TL-initial, primes in our form-priming experiments—i.e., Perea & Lupker, 2003.) The aim of Experiment 3 was to replicate the main findings of Experiments 1 and 2 in a single experiment. Finally, Experiment 4 was designed to provide a further examination of whether there are semantic/associative priming effects when TL-final nonword primes are used.

EXPERIMENT 1

As was noted above, Bourassa and Besner (1998) reported significant semantic/associative priming, using RL nonwords as primes. This effect, however, was quite small (7 msec) and required a rather large number of participants in order for it to be significant. In fact, in general, masked semantic/associative priming effects do not tend to be very large (around 10–26 msec across studies). As has been demonstrated by Pexman and colleagues (Pexman & Lupker, 1999; Pexman, Lupker, & Jared, 2001; see also Joordens & Becker, 1997, and Stone & Van Orden, 1993), however, word latencies and effect sizes in lexical decision tasks tend to be larger when pseudohomophone nonwords (e.g., *cleen*) are used. Therefore, in Experiment 1, we also manipulated the nature of the nonwords. Half of the participants received the more standard pronounceable nonwords (e.g., *gleek*), and the other half received pseudohomophones, in an attempt to provide a more sensitive examination of these effects.

Method

Participants. A total of 120 University of Western Ontario undergraduate students served as participants for course credit. All had normal or corrected-to-normal vision and were native speakers of English.

Materials. The targets were 120 words four to six letters in length (mean word frequency per one million words in the Kučera & Francis, 1967, count, 105; range, 1–1,600; mean Coltheart's *N*, 4.7; range, 0–16; mean number of letters, 4.88). (Coltheart's *N* is defined as the number of words differing by a single letter from the stimulus while preserving letter positions; e.g., *court* has only one "neighbor," *count*, so its *N* index is 1; Coltheart, Davelaar, Jonasson, & Besner, 1977.) The targets were presented in uppercase and were preceded by primes in lowercase that were (1) words associated to the target (related word condition; e.g., *judge*—*COURT*), (2) TL nonwords created by transposing two internal letters from the associated prime (the related TL-internal condition; e.g., *jugde*—*COURT*), (3) RL nonwords created by replacing an interior letter from the associated prime (related RL-nonword condition; e.g., *judpe*—*COURT*), (4) unrelated words (unrelated

word condition; e.g., *never*—*COURT*), (5) unrelated TL-internal nonwords (unrelated TL-internal condition; e.g., *neevr*—*COURT*), or (6) unrelated RL nonwords (unrelated RL-nonword condition; e.g., *nemer*—*COURT*). Word primes, TL-internal primes, and RL-nonword primes were rotated throughout the related and unrelated conditions so that each target word was primed by each of the six types of primes across the experiment. Thus, six sets of materials were constructed so that each target word would appear once in each set, but each time in a different priming condition. Different groups of participants ($n = 20$) were used for each set of materials. In all cases, the primes had five letters.² The related pairs are given in the Appendix.

Two additional sets of 120 nonwords four to six letters in length were selected. The standard nonwords were created by replacing a letter from a real English word (mean Coltheart's *N*, 3.7; range, 1–17; mean number of letters, 4.88), whereas the pseudohomophones (mean Coltheart's *N*, 4.1; range, 0–20; mean number of letters, 4.77) were selected from previous lexical decision experiments in which pseudohomophones were used (Pexman & Lupker, 1999; Pexman et al., 2001). The 120 nonword targets were preceded by 40 unrelated word primes and 80 unrelated nonword primes. Half of the participants received the standard pronounceable nonwords, and the other half received pseudohomophones. The nonwords used in the standard nonword and pseudohomophone conditions are also listed in the Appendix.

Procedure. The participants were run individually in a sound-attenuated room. Each trial consisted of a sequence of four visual events. The first was a forward mask consisting of a row of five hash marks (#####). This mask was presented for 500 msec. The mask was immediately followed by the prime in lowercase letters exposed for 40 msec, which was in turn immediately followed by a row of five hash marks (#####) for 40 msec. Finally, the target in uppercase letters replaced the mask and remained on the screen until the response. (This procedure was the same as that used by Bourassa & Besner, 1998. Note that it allows for a longer stimulus onset asynchrony [SOA; 80 msec] without increasing the duration of the prime. As such, it may increase the chances of obtaining masked semantic/associative priming effects, effects that are difficult to obtain with SOAs of 50 msec or less; see Forster et al., 2003.) Each stimulus was centered in the viewing screen and, hence, occupied the same position as did the preceding stimulus.

The items were presented on a TTX Multiscan Monitor (Model 3435P). Presentation was controlled by a Trillium Computer Resources PC. Words appeared as white characters on a black background. Reaction times (RTs) were measured from target onset until the participant's response. The participants were asked to classify the letter sequence presented in uppercase letters as a word or a nonword. No mention was made of the number of stimuli that would be presented on each trial. The participants indicated their decisions by pressing one of two response buttons. When the participant responded, the target disappeared from the screen. Each participant received a different pseudorandom ordering of items. Each participant also received 20 practice trials (with the same manipulation as that in the experimental trials) prior to the 240 experimental trials. The whole session lasted approximately 16 min.

Results and Discussion

Incorrect responses (2.1% of the data for word targets) and RTs less than 250 msec or greater than 1,200 msec (less than 2.7% of the data for word targets) were excluded from the latency analysis. Mean latencies for correct responses and error rates were calculated across individuals and items.³ Subjects and items analyses of variance (ANOVAs) based on the participants' response latencies and percentages of errors in each block were conducted on the basis of a 2 (relatedness: related or unrelated) $\times$ 3 (type of

prime: word prime, TL-internal nonword prime, or RL-nonword prime) $\times$ 2 (type of nonword: standard nonwords or pseudohomophones) $\times$ 6 (list: List 1, List 2, List 3, List 4, List 5, or List 6) design. The list factor was included as a dummy variable to extract the variance due to the error associated with the lists (see Pollatsek & Well, 1995). The mean RTs and error percentages from the subjects analyses are presented in Table 1.

RT analyses. Not surprisingly, the main effect of relatedness was significant [$F_1(1,108) = 13.13, MS_e = 1,218.7, p < .001; F_2(1,114) = 12.77, MS_e = 2,427.3, p < .001$], indicating that we did observe a semantic/associative priming effect. Also significant were the main effects of type of nonword [$F_1(1,108) = 9.68, MS_e = 51,321.3, p < .003; F_2(1,114) = 99.28, MS_e = 7,712.9, p < .001$] and type of prime [$F_1(2,216) = 4.68, MS_e = 1,144.2, p < .001; F_2(2,228) = 4.18, MS_e = 2,049.1, p < .02$]. These results reflect the fact that word target latencies were shorter (1) with standard nonwords and (2) when the (masked) primes were words. The interaction between relatedness and type of prime was marginally significant [$F_1(2,216) = 2.71, MS_e = 899.9, p < .07; F_2(2,228) = 2.62, MS_e = 21,999.9, p < .08$]. Planned comparisons, however, revealed that the semantic/associative priming effect was significant only for word primes [14.5 msec; $F_1(1,108) = 11.50, MS_e = 1,144.6, p < .001; F_2(1,114) = 11.74, MS_e = 2,457.2, p < .002$] and for TL-internal nonword primes [11 msec; $F_1(1,108) = 6.97, MS_e = 1,055.6, p < .011; F_2(1,114) = 8.27, MS_e = 1,997.5, p < .006$], and not for the RL-nonword primes (3 msec; both F s < 1). The other interactions did not approach significance (all $ps > .15$).

Error analyses. The main effect of type of nonword was significant [$F_1(1,108) = 6.12, MS_e = 37.20, p < .015; F_2(1,114) = 11.74, MS_e = 38.81, p < .001$]. The participants made fewer errors (to the word targets) when the nonwords were pseudohomophones than when they were standard nonwords. Although the main effect of relatedness was very small, it was marginally significant [$F_1(1,108) = 3.68, MS_e = 10.27, p = .057; F_2(1,114) = 2.82, MS_e =$

26.78, $p < .10$]. The error rate to word targets was 0.4% lower following related primes (2.2%) than following unrelated primes (2.6%). Although the interaction between relatedness and type of prime was not significant (both $ps > .15$), a subsequent analysis showed that there was a relatedness effect only for the word primes [1.8% vs. 2.7% for the related and the unrelated conditions, respectively; $F_1(1,108) = 5.04, MS_e = 10.94, p < .03; F_2(1,114) = 5.41, MS_e = 20.36, p < .002$], and not for the TL-internal nonword primes (2.2% vs. 2.6%) or the RL-nonword primes (2.6% vs. 2.6%; all $ps > .15$).

The results were clear. There was a semantic/associative priming effect with both word primes and TL-internal nonword primes (14.5 and 11 msec, respectively), but not with RL-nonword primes (3 msec).⁴ Thus, TL-internal nonword primes do appear to access semantic/associative information from their base words, producing priming of related words, whereas the evidence that RL-nonword primes do so as well is quite limited.

A point that should be noted here is that the differences in the sizes of the priming effects for the two types of nonword primes were due mainly to the differences in the unrelated, control conditions. That is, there was significant priming in the TL-internal condition, but not in the RL-nonword condition, not because target latencies were shorter in the related TL-internal condition than in the related RL-nonword condition, but because target latencies were longer in the TL-internal control condition than in the RL-nonword control condition.

In both cases, the unrelated, control conditions were created by re-pairing the primes and the targets from the parallel related conditions. Thus, the primes and the targets were identical in the related and the unrelated conditions in both the TL-internal and the RL-nonword conditions, creating the most appropriate control conditions. The question remains, however, as to whether this particular aspect of the data provides a challenge for our claim that TL-internal nonword primes produce semantic/associative priming, whereas, if they do so at all, RL-nonword primes do so to a much lesser degree.

In considering this issue, we first will focus on the RL-nonword prime condition. The manipulation in this condition, in which standard nonwords were used as foils, was a virtual replication of Bourassa and Besner's (1998) manipulation (i.e., we used virtually all of their stimulus pairs, along with a few added ones). The results were also a virtual replication. The size of our priming effect for the RL-nonword primes when standard nonword foils were used (6 msec) was virtually identical to their 7-msec effect. Thus, it seems rather unlikely that the contrast between our related condition and our unrelated, control condition when RL-nonword primes were used (i.e., the size of our priming effect) provides a misleading picture. As Bourassa and Besner have argued and as the results of Experiment 1 do not contradict, RL nonwords may be able to produce very small semantic/associative priming effects.

We will focus next on the claim that TL-internal nonword primes do produce semantic/associative priming. Is

Table 1
Mean Response Times (RTs, in Milliseconds) and Percentages of Errors for Word Targets in Experiment 1

Type of Prime	Condition					
	Related		Control		Priming	
	RT	%	RT	%	RT	%
Standard nonwords						
Word primes	633	1.5	647	2.1	14	0.6
TL-internal primes	644	1.7	656	1.8	12	0.1
RL-nonword primes	645	2.3	651	1.8	6	-0.5
Pseudohomophones						
Word primes	686	2.0	701	3.3	15	1.3
TL-internal primes	696	2.7	706	3.4	10	0.7
RL-nonword primes	701	3.0	701	3.4	0	0.4

Note—The error rate for the nonwords was 7.5% for both standard nonwords and pseudohomophones; the mean correct response times for the nonwords were 808 and 823 msec for the standard nonwords and the pseudohomophones, respectively. TL, transposed letter; RL, replacement letter.

this claim being supported because the unrelated control condition, even though it was the most appropriate control condition, produced a spuriously long mean latency? The basic difficulty in evaluating this question is that little, if anything, is known about what aspects of a prime might affect the latency of an unrelated target, particularly, in a masked priming situation. Thus, it is simply impossible to predict, a priori, what the relative difficulties of the different control conditions should be. As such, given our present state of knowledge, the performance level in the unrelated control condition can be determined only empirically.

When we take an empirical perspective, our previous work (Perea & Lupker, 2003) clearly indicates that different prime types do produce differences in mean latencies for unrelated targets of approximately the sizes observed here. For example, in Experiment 2 of Perea and Lupker, unrelated TL-internal nonword primes led to target latencies that were 5 msec longer than target latencies following unrelated TL-final nonword primes. In Experiment 3, in which a completely different set of primes and targets was used, the 5-msec difference was reversed and, in addition, target latencies following unrelated word primes were 9 msec shorter than those following unrelated TL-internal nonword primes. We also know, of course, that different types of unrelated (and neutral) primes can have quite strong effects on target latencies in unmasked priming situations (e.g., de Groot, Thomassen, & Hudson, 1982; Ratcliff & McKoon, 1995; Zeelenberg, Pecher, De Kok, & Raaijmakers, 1998). Thus, although there is a dearth of knowledge about what aspects of unrelated (or neutral) primes might produce differences in target latencies, there is no a priori reason to believe that the fact that the mean latency in the TL-internal control condition was 5 msec longer than that in the RL-nonword control condition and 7 msec longer than that in the word prime control condition was a spurious result.

The real question in this set of experiments, however, is what the patterns of priming effects are for TL-internal nonword, RL-nonword, and, in subsequent experiments, TL-final nonword primes. These are, at their core, empirical questions. If, indeed, the results for TL-internal nonword primes in Experiment 1 overestimated the priming effect or the results for RL-nonword primes in Experiment 1 underestimated the priming effect, that fact should become apparent on the basis of the pattern of results in our subsequent experiments.

Note finally that the use of pseudohomophones did increase the response times and the error rates for word targets but did not increase the sizes of the priming effects. Thus, it is possible that priming effects are somewhat different than the types of effects that Pexman and colleagues (Pexman & Lupker, 1999; Pexman et al., 2001) were investigating (e.g., homophone effects and polysemy effects). Consistent with this idea, in a recent study, Drieghe and Brysbaert (2002) also found that masked semantic/associative priming effects did not increase in size with pseudohomophones versus standard nonwords. In con-

trast, Joordens and Becker (1997) did report larger semantic priming effects with pseudohomophones than with standard nonwords in their Experiment 2, although not in their Experiment 1.

EXPERIMENT 2

In Experiment 2, the question was whether we would obtain a semantic/associative priming effect for TL-nonword primes that were created by transposing letters in the final two positions (e.g., *judeg*–COURT). As was suggested earlier, position information about end letters may be coded more accurately than position information about internal letters. Thus, it is possible that the orthographic representation for *judeg* is somewhat less similar to that for *JUDGE* than is the orthographic representation for *judge*, implying that there would be very little evidence of semantic/associative priming from TL-final nonword primes. Also included, for purposes of comparison and replication, was the RL-nonword condition. Given that the pseudohomophone manipulation did not alter the size of the priming effects in Experiment 1, we decided to use only standard nonwords in Experiment 2.

Method

Participants. A total of 36 University of Western Ontario undergraduate students participated for course credit. All had normal or corrected-to-normal vision and were native speakers of English. None of these individuals had participated in Experiment 1.

Materials. The 120 target words and the 120 target (standard) nonwords were the same as those in Experiment 1. The targets were presented in uppercase and were preceded by primes in lowercase that were (1) words associated to the target (related word condition; e.g., *judge*–COURT), (2) TL nonwords created by transposing the two final letters from the associated prime (related TL-final condition; e.g., *judeg*–COURT), (3) replacement-letter nonwords created by replacing the fourth or fifth letters from the associated prime (related RL-nonword condition; e.g., *judpe*–COURT), (4) unrelated words (unrelated word condition; e.g., *never*–COURT), (5) unrelated TL-final nonwords (unrelated TL-final condition; e.g., *nevre*–COURT), or (6) unrelated RL nonwords (unrelated RL-nonword condition; e.g., *nevem*–COURT). In all cases, the primes had five letters. As in Experiment 1, six sets of materials were constructed so that each target word appeared once in each set, but each time in a different priming condition. Different groups of participants were used for each set of materials. The nonword targets were the same as those in the standard nonword condition in Experiment 1.

Procedure. The procedure was the same as that in Experiment 1.

Results and Discussion

Incorrect responses (4.6% of the data for word targets) and RTs less than 250 msec or greater than 1,200 msec (less than 1.2% of the data for word targets) were excluded from the latency analysis. Mean latencies for correct responses and error rates were calculated across individuals and items. Subjects and items ANOVAs based on the participants' response latencies and percentages of errors in each block were conducted on the basis of a 2 (relatedness: related or unrelated) $\times$ 3 (type of prime: word prime, TL-final nonword prime, or RL-nonword prime) $\times$ 6 (list: List 1, List 2, List 3, List 4, List 5, or List 6) design. The

Table 2
Mean Response Times (RTs, in Milliseconds) and
Percentages of Errors for Word Targets in Experiment 2

Type of Prime	Condition					
	Related		Control		Priming	
	RT	%	RT	%	RT	%
Word primes	556	3.9	585	6.1	29	2.1
TL-final primes	569	3.3	574	4.3	5	1.0
RL-nonword primes	579	4.2	583	5.8	4	1.6

Note—The mean correct response times and error rates for nonwords were 678 msec and 8.5%. TL, transposed letter; RL, replacement letter.

mean RTs and percentages of errors from the subjects analyses are presented in Table 2.

RT analyses. The main effect of relatedness was significant [$F_1(1,30) = 12.44$, $MS_e = 653.5$, $p < .001$; $F_2(1,114) = 16.64$, $MS_e = 2,799.9$, $p < .001$]. Thus, once again, we observed a significant semantic/associative priming effect. Also, again the main effect of type of prime was significant [$F(2,60) = 3.39$, $MS_e = 723.4$, $p < .05$; $F_2(2,228) = 2.87$, $MS_e = 3,043.9$, $p = .059$]. Target latencies were shorter following word primes. Both of these effects, however, were qualified by the significant interaction between relatedness and type of prime [$F(2,60) = 5.94$, $MS_e = 622.1$, $p < .005$; $F_2(2,228) = 5.07$, $MS_e = 2,570.7$, $p < .008$]. Planned comparisons showed that the semantic/associative priming effect was significant for word primes [29 msec; $F_1(1,30) = 38.23$, $MS_e = 389.9$, $p < .001$; $F_2(1,114) = 20.70$, $MS_e = 2899.3$, $p < .001$], but not for TL-final nonword primes (5 msec) or for RL-nonword primes (4 msec; all $F_s < 1$).

Error analyses. The ANOVA on the error data indicated only a main effect of relatedness [$F_1(1,30) = 5.70$, $MS_e = 24.88$, $p < .025$; $F_2(1,114) = 5.64$, $MS_e = 83.85$, $p < .02$]. The error rate to word targets was 1.6% lower following a related prime (3.8%) than following an unrelated prime (5.4%). This effect did not interact with type of prime (both $p_s > .15$); however, a subsequent analysis indicated that the priming effect in the error data essentially occurred only for the word primes [$F_1(1,30) = 4.35$, $MS_e = 20.42$, $p < .05$; $F_2(1,114) = 4.85$, $MS_e = 61.05$, $p < .04$], and not for the TL-final nonword primes or the RL-nonword primes (both $p_s > .15$).

Again, the results were clear. There was a semantic/associative priming effect for word primes, but there was only minimal evidence of a priming effect for TL-final nonword primes or RL nonword primes. Thus, it appears that neither RL nonwords nor TL nonwords created by exchanging the final two letters are particularly effective at activating semantic/associative information from their base words.

One might also note that, in line with the preceding discussion about the differences between control conditions, the unrelated TL-final control condition produced target latencies that were 11 msec faster than those for the unrelated word prime condition and 9 msec faster than those for the unrelated RL-nonword control condition. Again, although there is no obvious explanation for these differ-

ences, there is also no reason to suggest that these means did not accurately reflect the processing difficulty of the respective control conditions.

EXPERIMENT 3

The results of Experiments 1 and 2 indicate that TL-internal nonwords do activate the semantic/associative representations of their base words, whereas TL-final nonwords do not (or at least, they are much less effective at doing so). One possible concern about this conclusion is that the TL-internal and the TL-final conditions were not contrasted directly with one another (i.e., the former was contained in Experiment 1, whereas the latter was contained in Experiment 2). Interestingly, one clear difference between the two experiments was that latencies were about 70 msec shorter in Experiment 2 than in Experiment 1 (even when only the standard nonword foil condition is considered). One could, therefore, attempt to explain the lack of a priming effect in the TL-final condition (in contrast to the clear priming effect in the TL-internal condition) merely by saying that TL priming effects are smaller when participants are fast. In an attempt to examine this possibility, Experiment 3 was designed to directly contrast the associative/semantic priming effects obtained with TL-internal and TL-final nonword primes with the same set of participants. For reasons of design efficiency, we focused on the critical conditions (the TL-internal condition vs. the TL-final condition), removing the conditions that involved RL-nonword primes and word primes.

Method

Participants. A total of 36 University of Western Ontario undergraduate students participated for course credit. All had normal or corrected-to-normal vision and were native speakers of English. None had participated in the previous experiments.

Materials. The 120 target words and the 120 target (standard) nonwords were the same as those in Experiments 1 and 2. The targets were presented in uppercase and were preceded by primes in lowercase that were (1) TL nonwords created by transposing two internal letters from the associated prime (the related TL-internal condition; e.g., *judge*—COURT), (2) TL nonwords created by transposing the two final letters from the associated prime (related TL-final condition; e.g., *judge*—COURT), (3) unrelated TL-internal nonwords (unrelated TL-internal condition; e.g., *neevr*—COURT), or (4) unrelated TL-final nonwords (unrelated TL-final condition; e.g., *nevre*—COURT). Four sets of materials were constructed so that each target word appeared once in each set, but each time in a different priming condition. Different groups of participants were used for each set of materials.

Procedure. The procedure was the same as that in Experiments 1 and 2.

Results and Discussion

Incorrect responses (2.5% of the data for word targets) and RTs less than 250 msec or greater than 1,200 msec (less than 0.7% of the data for word targets) were excluded from the latency analysis. Mean latencies for correct responses and error rates were calculated across individuals and items. Subjects and items ANOVAs based on the par-

Table 3
Mean Response Times (RTs, in Milliseconds)
and Percentage of Errors for Word Targets in Experiment 3

Type of Prime	Condition					
	Related		Control		Priming	
	RT	%	RT	%	RT	%
TL-internal primes	556	1.9	571	3.0	15	1.1
TL-final primes	559	3.1	562	2.0	3	-1.1

Note—The mean correct response times and error rates for nonwords were 654 msec and 4.3%. TL, transposed letter.

ticipants' response latencies and percentages of errors in each block were conducted on the basis of a 2 (relatedness: related or unrelated) $\times$ 2 (type of prime: TL-internal prime or TL-final prime) $\times$ 4 (List: List 1, List 2, List 3, or List 4) design. The mean RTs and percentages of errors from the subjects analyses are presented in Table 3.

RT analyses. The main effect of relatedness was significant [$F_1(1,32) = 9.64$, $MS_e = 302.8$, $p < .005$; $F_2(1,116) = 10.28$, $MS_e = 1,061.6$, $p < .003$]; responses to target words were 9 msec faster when these target words were preceded by TL primes than when they were preceded by control primes. The main effect of type of prime was not significant [$F_1(1,32) = 1.42$, $MS_e = 275.0$, $p > .15$; $F_2(1,116) = 0.45$, $MS_e = 1,474.8$, $p > .15$]. More important, the effect of relatedness was qualified by the significant interaction between relatedness and type of prime [$F_1(1,32) = 4.93$, $MS_e = 255.1$, $p < .035$; $F_2(1,116) = 4.13$, $MS_e = 1,329.6$, $p < .045$]. Planned comparisons showed that the semantic/associative priming effect was significant for TL-internal nonword primes [15 msec; $F_1(1,32) = 12.59$, $MS_e = 318.1$, $p < .002$; $F_2(1,116) = 12.84$, $MS_e = 1,242.2$, $p < .001$], but not for TL-final nonword primes (3 msec; both F s < 1).

Error analyses. The ANOVA on the error data produced a significant interaction only between relatedness and type of prime [$F_1(1,32) = 5.04$, $MS_e = 7.41$, $p < .035$; $F_2(1,116) = 4.90$, $MS_e = 25.38$, $p < .03$]. The error rate to word targets was 1.1% lower following a related TL-internal nonword prime (1.9%) than following an unrelated TL-internal nonword prime (3.0%), whereas the error rate to word targets was 1.1% higher following a related TL-final nonword prime (3.1%) than following an unrelated TL-final nonword prime (2.0%).

Again, the results were clear. Reinforcing the results of Experiment 1, there was a significant semantic/associative priming effect for TL-internal nonword primes (15 msec). Reinforcing the results of Experiment 2, there was only a small and nonsignificant priming effect for TL-final nonword primes (3 msec). Thus, as was argued previously, it appears that TL nonwords created by exchanging two internal letters are much more effective at activating semantic/associative information from their base words than are TL nonwords created by exchanging the two final letters.

Interestingly, paralleling Experiment 1, the differential priming effects were due mainly to differences in the control conditions. That is, the unrelated TL-final control condition produced shorter target latencies than did the

unrelated TL-final control condition. As a result, although the two related conditions produced similar target latencies, only the TL-internal condition produced a significant priming effect.

EXPERIMENT 4

The results of Experiments 1 and 3 have demonstrated that TL-internal nonwords activate the semantic/associative representations of their base words. In contrast, TL-final nonwords appear to be much less effective at doing so (nonsignificant effects of 5 and 3 msec in Experiments 2 and 3). Once again, however, there might be some concern that we have underestimated the size of the TL-final priming effect due to the fact that, if we look at the related conditions of Experiment 3 (*jugde* vs. *judge* as primes), there is little difference between the TL-internal and the TL-final conditions.

For that reason, we decided to look once again for the presence of associative/semantic priming effects with TL-final nonword primes. In Experiment 4, for reasons of design efficiency, we focused just on the two critical conditions: the related TL-final condition versus the unrelated TL-final control condition. There were, therefore, 60 prime-target pairs in each condition for each participant.

Method

Participants. A total of 18 University of Western Ontario undergraduate students participated for course credit. All had normal or corrected-to-normal vision and were native speakers of English. None had participated in the previous experiments.

Materials. The 120 target words and the 120 target (standard) nonwords were the same as those in Experiments 1–3. The targets were presented in uppercase and were preceded by primes in lowercase that were (1) TL nonwords created by transposing the two final letters from the associated prime (related TL-final condition; e.g., *judge*–*COURT*) or (2) unrelated TL-final nonwords (unrelated TL-final condition; e.g., *nevre*–*COURT*). Two sets of materials were constructed so that each target word appeared once in each set, but each time in a different priming condition. Different groups of participants were used for each set of materials.

Procedure. The procedure was the same as that in Experiments 1–3.

Results and Discussion

Incorrect responses (3.1% of the data for word targets) and RTs less than 250 msec or greater than 1,200 msec (less than 2.0% of the data for word targets) were excluded from the latency analysis. Mean latencies for correct responses and error rates were calculated across individuals and items. Subjects and items ANOVAs based on the par-

Table 4
Mean Response Times (RTs, in Milliseconds)
and Percentages of Errors for Word Targets in Experiment 4

Type of Prime	Condition					
	Related		Control		Priming	
	RT	%	RT	%	RT	%
TL-final primes	597	3.3	600	3.0	3	-0.3

Note—The mean correct response times and error rates for nonwords were 715 msec and 10.7%. TL, transposed letter.

ticipants' response latencies and percentages of errors in each block were conducted on the basis of a 2 (relatedness: related or unrelated) $\times$ 2 (list: List 1 or List 2) design. The mean RTs and percentages of errors from the subjects analyses are presented in Table 4.

There were no signs of a relatedness effect in the latency data (3 msec) or in the error data (-0.3% ; all $F_s < 1$).

The results were again clear. As in Experiments 2 and 3, TL nonwords created by transposing two final letters do not seem to have been very effective at activating semantic/associative information from their base words. Across the three experiments, the associative/semantic priming effects with TL-final nonwords were 5, 3, and 3 msec in Experiments 2, 3, and 4, respectively, with a combined analysis based on all three experiments also failing to show a reliable priming effect. Of course, one can never prove the null hypothesis, and the fact that, in all these experiments, the relatedness effect was positive suggests that TL-final nonwords may activate, to a small degree, the associative/semantic representation of their base word. If this effect is real, however, it is clearly very small.

GENERAL DISCUSSION

In terms of the main empirical goal of this research, the results clearly demonstrate that TL-internal nonword primes produce a reliable masked semantic/associative priming effect (around 10–15 msec, which was quite similar to the size of the priming effect for word primes in Experiment 1). In contrast, TL-final nonword primes and RL-nonword primes show only minimal evidence of being able to produce semantic/associative priming effects (around 3–5 msec in all cases). The greater size of the associative/semantic priming effects for TL-internal, in comparison with TL-final, primes is consistent with the finding that TL-internal primes yield greater form-priming effects than do TL-final primes (Perea & Lupker, 2003). The implication is that the orthographic representations of TL-internal nonwords are fairly similar to those of their base words, certainly much more similar than the orthographic representations of TL-final nonwords and RL nonwords are to their base words.

In trying to evaluate what the present results tell us about position coding, we first focus mainly on the TL-internal priming effects. One possible implication of the existence of priming from TL-internal, but not TL-final, nonwords is that although end letters may be rapidly tied to their positions within the orthographic representation, position coding for the internal letters takes time to develop (see Adams, 1979). That is, at least for internal letter positions, letter strings may be encoded more rapidly in terms of the identities of their individual letters than in terms of their absolute positions (see Humphreys et al., 1990). Therefore, the orthographic representation would contain the letters *d* and *g* when the letter string *judge* was processed (as well as when *judge* was processed) well before their positions had been determined. As such, it would be expected that the orthographic representation derived from

judge and that derived from *judge* would both activate the lexical representation for JUDGE, at least early in processing.

If one were willing to make this assumption, one could then explain the TL-internal priming results in terms of most models of visual word recognition (i.e., in terms of most position-coding mechanisms). For example, in terms of the channel-specific models, the argument would be that *judge* primes COURT because, until the letters are assigned to their channels, the system does not know that *judge* is not *judge* and, hence, the lexical representation for JUDGE is initially activated (as well as its semantic/associative representation). The problem this creates for this type of model, however, is that it then requires the model to have a second (slow-acting) mechanism for allocating already encoded letters to their positions—that is, a separate mechanism that ultimately determines that the *g* is in the third position and not in the fourth (or second) and that the *d* is in the fourth position and not in the third. At present, none of these models incorporates such a mechanism.

One possible way to examine the plausibility of such a two-stage model would be to use a signal-to-respond paradigm (Doshier, 1976; Ratcliff & McKoon, 1982; Reed, 1973; see also Hintzman & Curran, 1997, for an application to the lexical decision task). In this technique, participants have to make a lexical decision at specific times, neither before nor (much) after the signal is presented. If the two-stage model is correct, when the lag between the stimulus and the signal to respond is very brief (e.g., 100 or 200 msec), the information about (internal) letter positions will not have been fully processed. This implies that participants should make more false alarms to TL nonwords (e.g., *judge*) than to RL nonwords (e.g., *judpe*) under these conditions. This difference should decrease/vanish when the lag is relatively long (e.g., 500 msec), since the incoming information concerning letter position would make the TL nonwords less similar to their base words. In agreement with this hypothesis, Gómez, Perea, and Ratcliff (2002) found that TL nonwords produced substantially more false positives than did RL nonwords at short lags (100, 200, and 300 msec), but not at long lags (500 and 1,000 msec).

It is worth noting that the proposal that letter identity information and letter position information might become separated in the perceptual system is not new (Estes, 1975; Ratcliff, 1981; Rumelhart & McClelland, 1982). Furthermore, one could easily implement this idea within the interactive-activation model by merely assuming that “information presented on one location might activate detectors in a range of locations rather than in one fixed position” (Rumelhart & McClelland, 1982, p. 89). In this way, upon presentation of the TL nonword *judge*, the letters *g* and *d* would activate the detectors at neighboring letter positions, and thereby, the lexical representation of the base word JUDGE would be activated to some extent.

In fact, empirical evidence suggests that the representation of a letter may initially be distributed over letter po-

sitions (Ratcliff, 1981). By the end of processing, the correct letter position for each letter would, of course, be determined due, at least in part, to feedback from higher level processes. Interestingly, this modified interactive-activation model would predict that the processing of TL words, such as *trial*, would suffer due to interference from the activation of their TL mates (e.g., *trail*), an effect that does seem to occur (see Andrews, 1996). In any case, if we do assume that assigning positions to letter identities takes some time, the masked nonword *judge* would be expected to activate the lexical representation of JUDGE to a reasonable extent, which would explain the presence of reliable semantic/associative effects with TL primes. At present, however, these notions have not been implemented in an interactive-activation type model.

Although it is possible that an implemented version of a modified interactive-activation model (or some other channel-specific model) might be able to explain the present results, there are a number of other problems for any models employing a simple *channel-specific* coding scheme. For instance, as McClelland (1986) acknowledged, the interactive-activation model had to include the unlikely assumption that there is a separate set of letter detectors for each letter position. Furthermore, in the original interactive-activation model, words with different numbers of letters do not activate each other's lexical representations (see Jacobs et al., 1998). That is, the word *hose* activates only the lexical representations for (formally similar) four-letter words, and therefore, it would not activate the lexical representation for HOUSE. Such a prediction is inconsistent with the available behavioral evidence (e.g., de Moor & Brysbaert, 2000; McClelland, 1986; Perea & Carreiras, 1998; Peressotti & Grainger, 1999).

Recently, two extensions of the interactive-activation model have been proposed that use a relative-position, rather than an absolute-position, coding scheme. These models do allow a word to activate the lexical representations of words with different lengths (multiple read-out model [MROM-p], Jacobs et al., 1998; dual-route cascaded [DRC] model, Coltheart et al., 2001). In the MROM-p, the external letters are used as anchor points, and the other letters in a word are represented in terms of their relative position in the word. For instance, the word *judge* would be encoded as *j* in the initial position, *u* in the initial plus one, *d* in the initial plus two, *g* in the final minus one, and *e* in the final position. The DRC model uses a single system for words of various lengths (up to eight letters), in which the first letter is used as an anchor point. For instance, the word *judge* would be encoded as *j* in the initial position, *u* in the second position, *d* in the third position, *g* in the fourth position, *e* in the fifth position, whereas in positions six to eight there would be a blank-letter character. The coding schemes in these computational models appear to be something of an improvement over the coding scheme of the original interactive-activation model; however, as they stand, neither of these coding schemes would predict that a TL nonword such as *judge* would activate the lexical representation of the base word JUDGE to even

the same degree as the RL nonword *judpe*. If we did assume that letter identity information and letter position information become separated in the perceptual system, as was discussed above, these models might, in principle, capture the observed effects. Simulations on such an implemented model would be necessary in order to verify this claim.

Another potential way to explain TL effects would be to use a parallel distributed processing model having a coarse coding scheme (see, e.g., McClelland, 1986; Mozer, 1987). In the PABLO model (McClelland, 1986), each detector serves only as a partial specification of a letter combination. For instance, the word *back* would be coded as ($_B Bx$) ($x A Ax$) ($x C Cx$) ($x K K_$), where the first and the last letters of the strings are used as anchor points (the $_$ sign indicates a blank, and the letter x indicates any letter). Simulations with the PABLO model show that the TL nonword *bcak* activates the connections for its base word, BACK, as much as *back* itself does (McClelland, 1986). (An RL nonword such as *bick* would activate BACK to a much lesser degree.) Because of the important role played by external letter positions in the model, TL similarity effects would be more likely when the transposition involved two internal letters. Thus, the coding scheme used in the PABLO model might allow that model to capture the effects found in the present experiments. (The implemented version of PABLO was somewhat limited, however. For instance, its alphabet contained only eight letters plus the “ $_$ ” sign, and its lexicon was composed of words from one to four letters in length.)⁵

Finally, another option that seems to capture TL similarity more naturally would be to use a model with a radically different coding scheme, such as that implemented in the SOLAR model (Davis, 1999). This model uses a spatial coding scheme in which letter codes are position independent. As a result, the TL nonword *judge* and its base word, JUDGE, share the same set of letter codes. To account for the fact that any orthographic coding scheme must ultimately be order sensitive, the order of letters in the SOLAR model is coded by the relative activity of the set of letter nodes. Thus, *judge* and *judge* share the same set of letter nodes, but they produce different activation patterns (e.g., in the word *judge*, the letter code corresponding to *j* is the one associated with the highest activation value, then the letter code corresponding to the letter *u* is associated with a slightly smaller activation value, and so on; see Figure 1). Because serial position is coded by relative activation, rather than by position-specific codes, and because of the way the network computes bottom-up input, *judge* and *judge* are more similar and, hence, more confusable than *judge* and *judpe*, essentially, throughout processing (see Figure 1). More specifically, in the SOLAR model, the computed orthographic similarity between the word *judge* and the TL-internal nonword *judge* is .80, which is substantially higher than the orthographic similarity between the word *judge* and the RL nonword *judpe* (.71),⁶ although it is, of course, less than the orthographic similarity between *judge* and itself (1.00). Thus, upon presentation of

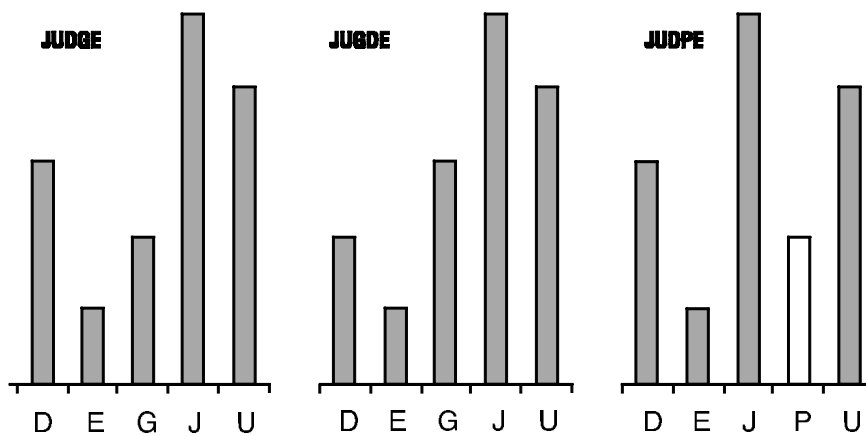


Figure 1. An example of spatial coding in the SOLAR model for the word *judge*, the transposed-letter nonword *jugde*, and the replacement-letter nonword *judpe*. The order in which letters occur is coded by the relative activity of nodes: The first letter is coded by the largest activation value, the second letter is associated with a slightly smaller activation value, and so on.

the masked TL nonword *jugde*, the expectation is that the lexical representation of the base word *JUDGE* would be at least partially activated (more so than from the RL nonword *judpe*), which could ultimately lead to a semantic/associative priming effect.

Note that the SOLAR model also gives a special role to the external letter positions (i.e., there is a weighting parameter in the model that favors a match in the end positions). The idea is that the external letters may be more easily coded than the internal letters, due to a lack of lateral inhibition (see also Whitney, 2001). As a result, nonwords created by transposing the two final letters would be less similar to their base words than are nonwords created by transposing two middle letters. Specifically, the orthographic similarity between the word *judge* and the TL-final nonword *judpe* in the model is only .74. This decrease in similarity between *jugde* and *judpe* is consistent with our failure to find a semantic/associative priming effect when TL-final nonwords were used as primes in Experiments 2–4. What should also be noted is that any model that hopes to simultaneously account for semantic/associative priming from TL-internal nonwords, but not from TL-final nonwords, will have to have a mechanism that gives a special status to the final letter.

CONCLUSIONS

The present research demonstrates the existence of masked semantic/associative priming effects produced by TL nonwords. Not surprisingly, there is more evidence that these effects are real when the letter transposition occurs in the middle of the word than when it occurs at the end of the word. These results are most consistent with models that propose that letters are position coded in a way that produces a high level of similarity between words and TL-internal nonwords (e.g., the SOLAR model) and are less consistent with models that assume some type of

channel-specific coding scheme, unless an additional set of assumptions is added to those models.

REFERENCES

- ADAMS, M. J. (1979). Models of word recognition. *Cognitive Psychology*, **11**, 133-176.
- ANDREWS, S. (1996). Lexical retrieval and selection processes: Effects of transposed-letter confusability. *Journal of Memory & Language*, **35**, 775-800.
- BODNER, G. E., & MASSON, M. E. J. (1997). Masked repetition priming for words and nonwords: Evidence for a nonlexical basis for priming. *Journal of Memory & Language*, **37**, 268-293.
- BODNER, G. E., & MASSON, M. E. J. (in press). Beyond spreading activation: An influence of relatedness proportion on masked semantic priming. *Psychonomic Bulletin & Review*.
- BOURASSA, D. C., & BESNER, D. (1998). When do nonwords activate semantics? Implications for models of visual word recognition. *Memory & Cognition*, **26**, 61-74.
- CHAMBERS, S. M. (1979). Letter and order information in lexical access. *Journal of Verbal Learning & Verbal Behavior*, **18**, 225-241.
- CLARK, H. H. (1973). The language-as-fixed-effect fallacy: A critique of language statistics in psychological research. *Journal of Verbal Learning & Verbal Behavior*, **12**, 335-359.
- COHEN, J. (1976). Random means random. *Journal of Verbal Learning & Verbal Behavior*, **15**, 261-262.
- 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.
- 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.
- DE GROOT, A. M. B., & NAS, G. L. J. (1991). Lexical representation of cognates and noncognates in compound bilinguals. *Journal of Memory & Language*, **30**, 90-123.
- DE GROOT, A. M. B., THOMASSEN, A. J. W. M., & HUDSON, P. T. W. (1982). Associative facilitation of word recognition as measured from a neutral prime. *Memory & Cognition*, **10**, 358-370.
- DE MOOR, W., & BRYSAERT, M. (2000). Neighborhood-frequency effects when primes and targets have different lengths. *Psychological Research*, **63**, 159-162.

- DOSHER, B. A. (1976). The retrieval of sentences from memory: A speed-accuracy study. *Cognitive Science*, **8**, 291-310.
- DRIEGHE, D., & BRYSAERT, M. (2002). Strategic effects in associative priming with words, homophones and pseudohomophones. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, **28**, 962-982.
- ESTES, W. K. (1975). The locus of inferential and perceptual processes in letter identification. *Journal of Experimental Psychology*, **104**, 122-145.
- ESTES, W. K., ALLMEYER, D. H., & REDER, S. M. (1976). Serial position functions for letter identification at brief and extended exposure durations. *Perception & Psychophysics*, **19**, 1-15.
- FORSTER, K. I. (1976). Accessing the mental lexicon. In R. J. Wales & E. W. Walker (Eds.), *New approaches to language mechanisms* (pp. 257-287). Amsterdam: North-Holland.
- 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*, **39A**, 211-251.
- FORSTER, K. I., MOHAN, K., & HECTOR, J. (2003). The mechanics of masked priming. In S. Kinoshita & S. J. Lupker (Eds.), *Masked priming: State of the art* (pp. 3-37). Hove, U.K.: Psychology Press.
- FRIEDMANN, N., & GVION, A. (2001). Letter position dyslexia. *Cognitive Neuropsychology*, **18**, 673-696.
- FROST, R., FORSTER, K. I., & DEUTSCH, A. (1997). What can we learn from the morphology of Hebrew? A masked priming investigation of morphological representation. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, **23**, 829-856.
- GÓMEZ, P., PEREA, M., & RATCLIFF, R. (2002, April). *Dinámica temporal de la activación léxica en pseudopalabras* [Time course of lexical activation in pseudowords]. Paper presented at the 4th meeting of the Sociedad Española de Psicología Experimental, Oviedo, Spain.
- GONNERMAN, L. M., & PLAUT, D. C. (2000, April). *Semantic and morphological effects in masked priming*. Poster presented at the 7th Annual Meeting of the Cognitive Neuroscience Society, San Francisco.
- GRAINGER, J. (1992). Orthographic neighborhoods and visual word recognition. In R. Frost & L. Katz (Eds.), *Orthography, phonology, morphology, and meaning* (pp. 131-146). Amsterdam: Elsevier.
- GRAINGER, J., & JACOBS, A. M. (1996). Orthographic processing in visual word recognition: A multiple read-out model. *Psychological Review*, **103**, 518-565.
- HINTZMAN, D. L., & CURRAN, T. (1997). Comparing retrieval dynamics in recognition memory and lexical decision. *Journal of Experimental Psychology: General*, **126**, 228-247.
- HOLMES, V. M., & NG, E. (1993). Word-specific knowledge, word-recognition strategies, and spelling ability. *Journal of Memory & Language*, **32**, 230-257.
- HUMPHREYS, G. W., EVETT, L. J., & QUINLAN, P. T. (1990). Orthographic processing in visual word recognition. *Cognitive Psychology*, **22**, 517-560.
- JACOBS, A. M., REY, A., ZIEGLER, J. C., & GRAINGER, J. (1998). MROM-p: An interactive activation, multiple readout model of orthographic and phonological processes in visual word recognition. In J. Grainger & A. M. Jacobs (Eds.), *Localist connectionist approaches to human cognition* (pp. 147-188). Mahwah, NJ: Erlbaum.
- JOHNSON, N. F., & PUGH, K. R. (1994). An examination of cohort models of visual word recognition: On the role of letters in word-level processing. *Cognitive Psychology*, **26**, 240-346.
- JOORDENS, S., & BECKER, S. (1997). The long and short of semantic priming effects in lexical decision. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, **23**, 1083-1105.
- JORDAN, T. R. (1990). Presenting words without interior letters: Superiority over single letters and influence of postmark boundaries. *Journal of Experimental Psychology: Human Perception & Performance*, **16**, 891-909.
- KEPPEL, G. (1976). Words as random variables. *Journal of Verbal Learning & Verbal Behavior*, **15**, 263-265.
- KUČERA, H., & FRANCIS, W. N. (1967). *Computational analysis of present-day American English*. Providence, RI: Brown University Press.
- LUKATELA, G., CARELLO, C., SAVIĆ, M., UROŠEVIĆ, Z., & TURVEY, M. T. (1998). When nonwords activate semantics better than words. *Cognition*, **68**, B31-B40.
- LUKATELA, G., & TURVEY, M. T. (1994). Visual lexical access is initially phonological: 1. Evidence from associative priming from words, homophones, and pseudohomophones. *Journal of Experimental Psychology: General*, **123**, 107-128.
- MASSON, M. E. J., & ISAAK, M. I. (1999). Masked priming of words and nonwords in a naming task: Further evidence for a nonlexical basis for priming. *Memory & Cognition*, **27**, 399-412.
- MCCLELLAND, J. L. (1986). The programmable blackboard model of reading. In J. L. McClelland & D. E. Rumelhart (Eds.), *Parallel distributed processing: Exploration in the microstructure of cognition. Vol. II: Psychological and biological models* (pp. 122-169). Cambridge, MA: MIT Press.
- MCCLELLAND, J. L., & RUMELHART, D. E. (1981). An interactive activation model of context effects in letter perception: Pt. 1. An account of basic findings. *Psychological Review*, **88**, 375-407.
- 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.
- NORRIS, D. (1986). Word recognition: Context effects without priming. *Cognition*, **22**, 93-361.
- 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. (1998). Orthographic neighbours are not all equal: Evidence using an identification technique. *Language & Cognitive Processes*, **13**, 77-90.
- PEREA, M., & CARREIRAS, M. (1998). Effects of syllable frequency and neighborhood syllable frequency in visual word recognition. *Journal of Experimental Psychology: Human Perception & Performance*, **24**, 1-11.
- PEREA, M., & GOTOR, A. (1997). Associative and semantic priming effects occur at very short SOAs in lexical decision and naming. *Cognition*, **67**, 223-240.
- PEREA, M., & LUPKER, S. J. (2003). 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, U.K.: Psychology Press.
- PEREA, M., & ROSA, E. (2002a). 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. (2000b). The effects of associative and semantic priming in the lexical decision task. *Psychological Research*, **66**, 180-194.
- PERESSOTTI, F., & GRAINGER, J. (1999). The role of letter identity and letter position in orthographic priming. *Perception & Psychophysics*, **61**, 691-706.
- PEXMAN, P. M., & LUPKER, S. J. (1999). The impact of semantic ambiguity on visual word recognition: Do homophone and polysemy effects co-occur? *Canadian Journal of Experimental Psychology*, **53**, 323-334.
- PEXMAN, P. M., LUPKER, S. J., & JARED, D. (2001). Homophone effects in lexical decision. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, **27**, 139-156.
- 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, & Cognition*, **21**, 785-794.
- RAAIJMAKERS, J. G. W., SCHRIJNEMAKERS, J. M. C., & GREMMEN, F. (1999). How to deal with "the language-as-fixed-effect fallacy": Common misconceptions and alternative solutions. *Journal of Memory & Language*, **41**, 416-429.

- RATCLIFF, R. (1981). A theory of order relations in perceptual matching. *Psychological Review*, **88**, 552-572.
- RATCLIFF, R., GÓMEZ, P., & MCKOON, G. (in press). A diffusion model account of the lexical decision task. *Psychological Review*.
- RATCLIFF, R., & MCKOON, G. (1982). Speed and accuracy in the processing of false statements about semantic information. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, **8**, 16-36.
- RATCLIFF, R., & MCKOON, G. (1995). Sequential effects in lexical decision: Tests of compound cue retrieval theory. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, **21**, 1380-1388.
- REED, A. V. (1973). Speed-accuracy trade-off in recognition memory. *Science*, **181**, 574-576.
- RUMELHART, D. E., & MCCLELLAND, J. L. (1982). An interactive activation model of context effects in letter perception: Pt. 2. The contextual enhancement effect and some tests and extensions of the model. *Psychological Review*, **89**, 60-94.
- SEIDENBERG, M. S., & MCCLELLAND, J. L. (1989). A distributed, developmental model of word recognition and naming. *Psychological Review*, **96**, 523-568.
- SERENO, J. A. (1991). Graphemic, associative, and syntactic priming effects at a brief stimulus onset asynchrony in lexical decision and naming. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, **17**, 459-477.
- SMITH, J. E. K. (1976). The assuming-will-make-it-so fallacy. *Journal of Verbal Learning & Verbal Behavior*, **15**, 262-263.
- STONE, G. O., & VAN ORDEN, G. C. (1993). Strategic control of processing in visual word recognition. *Journal of Experimental Psychology: Human Perception & Performance*, **19**, 744-774.
- TAFT, M., & VAN GRAAN, F. (1998). Lack of phonological mediation in a semantic categorization task. *Journal of Memory & Language*, **38**, 203-224.
- 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.
- WIKE, E. L., & CHURCH, J. D. (1976). Comments on Clark's "The language-as fixed-effect fallacy." *Journal of Verbal Learning & Verbal Behavior*, **15**, 249-255.
- WILLIAMS, J. N. (1996). Is automatic priming semantic? *European Journal of Cognitive Psychology*, **8**, 113-161.
- ZEELNBERG, R., PECHER, D., DE KOK, D., & RAAIJMAKERS, J. G. W. (1998). Inhibition for nonword primes in lexical decision re-examined: The influence of instructions. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, **24**, 1068-1079.

NOTES

1. Some visual word recognition models (e.g., the search model, Forster, 1976; the checking model, Norris, 1986) do remain neutral as to the specific coding scheme.

2. In the present experiments, we used 78 out of the 80 word-word pairs used by Bourassa and Besner (1998) plus 42 new pairs. (The two pairs we dropped from Bourassa and Besner's stimulus set were *spoke*-SPEAK and *bread*-WATER, which were replaced by *spoke*-TALK and *bread*-BUTTER.)

3. Although Clark (1973) has argued that items, as well as subjects, should be considered as a random factor in these types of analyses, the selection of items is seldom *random* in any sense of the term. Such is also the case in the present experiments. As such, as Wike and Church (1976) and others (e.g., Cohen, 1976; Keppel, 1976; Smith, 1976) have argued, item analyses would clearly be inappropriate in the present situation for a number of reasons, not the least of which is their strong negative bias (see also Raaijmakers, Schrijnemakers, & Gremmen, 1999). Nonetheless, for the interested reader, the results of item analyses will be reported. Conclusions, however, will be based only on the results from the subjects analyses. (Interested readers may also note that, in the present set of experiments, the results of the item analyses mimicked the results of the subjects analyses.)

4. It is worth noting that the leading edge of the group RT distributions (.1 quantile) for *all* correct responses (see Ratcliff, Gómez, & McKoon, in press) mimicked the analyses of mean RTs: The semantic/associative priming effect was 11 msec for word primes [532 vs. 543 msec; $F_1(1,108) = 7.93, MS_e = 867.3, p < .006$], 8 msec for TL-internal nonword primes [542 vs. 550 msec; $F_1(1,108) = 4.52, MS_e = 904.9, p < .04$], and only 2 msec for RL-nonword primes (544 vs. 546 msec; $F_1 < 1$).

5. In the connectionist model of Seidenberg and McClelland (1989), the coding scheme was based on triads of ordered letters, the so-called *wickelfeatures*. For instance, the codes for *judge* would be _ju, jud, udg, dge, ge_, where the _ sign refers to the end of the letter string. It seems unlikely that this type of coding scheme could explain TL similarity effects, because the TL nonword *judge* would be less similar to its base word than the RL nonword *judpe* (the wickelfeatures for these two nonwords are _ju, jug, ugd, gde, de_, and _ju, jud, udp, dpe, and pe_, respectively).

6. These values were obtained using the default parameters in Match-Calculator, an application written by Colin Davis to compute the orthographic similarity between two letter strings according to the SOLAR model. We would like to thank Colin Davis for providing us with this program and for corroborating our calculations (C. Davis, personal communication, February 25, 2003).

APPENDIX

Related Pairs in Experiments 1-4

The items are arranged in sextuplets in the following order: word prime, TL-internal prime, RL-nonword prime (Experiment 1), TL-final prime, RL-nonword prime (Experiment 2), target word.

never, neevr, nemer, nevre, nevem, ALWAYS; fruit, friut, freit, fruti, fruil, APPLE; march, macrh, manch, marhc, maroh, APRIL; Uncle, unlce, unvle, uncel, uncte, AUNT; roast, raost, roant, roats, roask, BEEF; above, abvoe, alove, aboev, abovs, BELOW; robin, roibn, rofin, robni, roban, BIRD; flesh, felsh, flosch, flehs, flest, BLOOD; skirt, skrit, skart, skitr, skist, BLOUSE; study, sutdy, stady, stuyd, stuky, BOOKS; house, huose, houge, houes, hounre, BRICK; carry, crary, caryg, caryr, carsy, BRING; hills, hlils, holls, hilsl, hille, BUMPS; bread, braed, breid, breda, breal, BUTTER; jails, jials, jaips, jaisl, jaits, CELLS; fraud, fruad, frald, fradu, fraod, CHEAT; miner, mienr, miver, minre, minen, COAL; judge, jugde, judpe, judeg, judpe, COURT; river, rievtr, ruver, rivre, rives, CREEK; thief, tihef, thirf, thife, thiaf, CROOK; faces, faecs, fapes, facse, facen, CROWD; fatal, faatl, fatil, fatla, fatac, DEATH; angel, agnel, antel, angle, angol, DEVIL; clean, claen, chean, clena, cleam, DIRTY; nurse, nusre, nunse, nures, nurss, DOCTOR; awake, awkae, awane, awaek, awaks, DREAM; lifts, litfs, lirts, liftst, lifte, DROPS; sober, soebr, siber, sobre, sobir, DRUNK; bacon, baocn, bamon, bacno, bacos, EGGS; knife, kinfe, knike, knief, FORK; empty, emtpy, empky, empyt, empky, FULL; plays, palys, pluys, plasy, plags, GAMES; sells, slsls, salls, selsl, selks, GOODS; coast, caost, coost, coats, coant, GUARD; curly, culry, cusly, curyl, curky, HAIR; glove, golve, glone, gloev, gloves, HAND; loves, loevs, lovos, lovse, lovis, HATES; spice, sipce, skice, spiec, spime, HERBS; mount, muont, moultr, moutn, mounk, HORSE; camel, caeml, cawel, camle, camek, HUMP; pearl, paerl, peirl, pearl, peard, JEWEL; royal, roayl, ropal, royla, royat, KINGS; early, ealry, eanly, earyl, earty, LATE; sneer, sener, skeer, snere, sneor, LAUGH; teach, taech, toach, teahc,

APPENDIX (Continued)

teash, LEARN; heavy, hevay, heamy, heayv, heamy, LIGHT; tiger, tiegr, tiper, tigre, tigem, LION; short, shrot, shart, shotr, shork, LONG; tight, tihgt, teght, tigh, tigt, LOOSE; maybe, mabye, marbe, mayeb, maybs, MIGHT; major, maojr, majar, majro, majos, MINOR; coins, cions, clins, coisn, coirs, MONEY; teeth, teteh, telth, teeht, teesh, MOUTH; usher, uhser, uther, ushre, ushen, MOVIE; bathe, bahte, buthe, bateh, batke, NAKED; scarf, scraf, scerf, scafr, scanf, NECK; green, geren, grenn, grene, greem, OLIVE; spray, spary, sroy, sprya, spre, PAINT; whole, whloe, wiole, whoel, whols, PARTS; paper, paepr, puper, papre, papec, PENCIL; pilot, pliot, pidot, pilito, pilok, PLANE; board, baord, boird, boadr, boart, PLANK; ideas, idaes, idoas, idesa, idean, PLANS; peach, paech, peash, peahc, peash, PLUM; songs, snogs, sengs, sonsg, sonps, POEMS; lakes, laeks, lokes, lakse, lakem, PONDS; lower, loewr, liwer, lowre, lowem, RAISE; shave, shvae, shawe, shaev, shane, RAZOR; coral, coarl, cural, corla, corak, REEFS; monks, mokns, manks, monsk, monke, ROBES; slide, silde, slike, slied, slidd, RULER; bolts, botls, bolps, bolst, boltn, SCREW; plant, palnt, plint, platn, plamt, SEEDS; point, piont, poilt, poitin, poind, SHARP; barns, banrs, balns, barsn, barrs, SHEDS; metal, meatl, megal, metla, metak, SHINY; pants, patns, palts, panst, pante, SHIRT; music, muisc, muric, musci, musoc, SING; cloud, cluod, choud, clodu, cloul, SKIES; large, lagre, lange, lareg, largs, SMALL; rough, ruogh, roegh, rouhg, rougl, SMOOTH; sleet, selet, sliet, sleed, SNOW; shoes, sheos, spoes, shose, shoen, SOCKS; couch, cuoch, cooch, couhc, courh, SOFA; lemon, leomn, leron, lemno, lemin, SOUR; north, nothr, norgh, norht, norlh, SOUTH; round, ruond, roond, roudn, rount, SQUARE; meats, maets, miats, meast, meaks, STEAK; rigid, riigd, rogid, rigdi, rigil, STIFF; stick, sitck, steck, stikc, stisk, STONE; bible, bilbe, beble, bibel, bibke, STORY; broom, borom, braom, bromo, broam, SWEEP; salty, satly, safty, salyt, salky, SWEET; chair, chiar, chuir, chari, chaim, TABLE; heads, haeds, hends, heasd, heaks, TAILS; given, giev, gilen, givne, givem, TAKEN; fairy, fiary, faimy, faiyr, faisy, TALES; speak, spaek, spes, speka, speaf, TALK; tells, tlels, tulls, telsl, telle, TALKS; onion, oinon, omion, onino, oniom, TEARS; fleas, flaes, fluas, fles, fleis, TICKS; clock, colck, cleck, clokc, cloct, TIME; train, trian, traln, trani, trais, TRACK; magic, maigc, magoc, magci, maguc, TRICK; stems, setms, stefs, stesm, stemm, TWIGS; angry, anrgy, angdy, angyr, angny, UPSET; coats, caots, clats, coast, coaks, VESTS; sleep, selep, slemp, slepe, sleip, WAKES; polka, pokla, ponka, polak, polke, WALTZ; needs, nedes, neeks, needd, neeks, WANTS; ocean, ocaen, oceln, ocena, oceam, WAVES; grass, garss, gnass, grass, grasm, WEEDS; wagon, waogn, wafon, wagno, wagos, WHEEL; black, balck, bleck, blak, blask, WHITE; corks, croks, corms, corsk, corls, WINES; girls, gilrs, girns, girsl, girks, WOMEN; sheep, shep, shoep, shepe, sheey, WOOL; globe, golbe, glibe, gloeb, gloke, WORLD; value, vaule, vanue, valeu, valce, WORTH; child, chlid, chuld, chidl, chil, YOUNG

Nonwords in Experiments 1–4

Standard nonwords: merse, glest, plust, drost, frent, bandom, nolk, chen, garrel, noast, plare, frope, cottle, plick, lurge, cleed, grire, vener, jore, quast, shir, codel, nurch, crich, reast, drave, merve, poose, chade, norke, roste, pruch, trisp, yinch, brape, proth, leath, glink, vind, gress, tream, tault, gripa, gifle, plich, falet, emple, trime, peash, gouch, reasy, telk, crompt, heast, arone, kelsh, bramp, fint, ferch, cheab, sholl, shoon, hald, nelect, mact, shrum, rooze, blass, benim, heech, naip, untle, grulp, calt, hurky, bline, pewor, hote, gleek, hoid, thark, wesp, blipe, bero, sloat, teason, fotion, appit, drick, trusk, wadge, frink, souch, thrag, fing, yownd, dake, dage, halst, crame, grosp, sooch, jaste, brich, vaste, chork, guilm, nulk, cremit, polt, treda, blatt, flish, parm, gour, bosh, zourk, kremp, woney, qual

Pseudohomophones (Experiment 1 only): phork, hoarn, scail, nues, fale, goast, creem, blaim, greef, roals, chace, spind, lern, rade, ceese, bloo, graive, fayze, crule, murge, soald, durt, cleen, braive, korts, neace, kerse, vurse, chuze, shair, rane, laff, perce, boal, koast, rute, spoart, werth, hazz, prufe, squair, reech, cheeze, sute, joaks, trupe, taist, dout, dait, daize, boarn, breth, heer, sope, doaps, swet, looce, nerce, brane, grone, falce, tode, deels, rong, fome, leep, mait, cort, phit, floar, frate, tutch, hed, dood, nifes, gane, gloab, wheal, raige, blud, sleap, plite, speek, smoak, scoar, sirch, kase, trax, stoar, meel, tipes, taim, shurt, hoze, bleek, hored, turse, wate, brief, cryde, jales, noize, gerls, teeche, voyce, thret, trale, grean, blurse, greese, pleez, whied, darck, fownd, woond, chrow, cheef, jooce, munny, coad