

Mechanisms of source confusion and discounting in short-term priming:

1. Effects of prime duration and prime recognition

DAVID E. HUBER

University of Colorado, Boulder, Colorado

RICHARD M. SHIFFRIN

Indiana University, Bloomington, Indiana

RAUSHANNA QUACH

University of Colorado, Boulder, Colorado

and

KEITH B. LYLE

Indiana University, Bloomington, Indiana

Huber, Shriffrin, Lyle, and Ruys (2001) measured short-term repetition priming effects in perceptual identification with two-alternative forced-choice testing. There was a preference to choose repeated words following passive viewing of primes and a preference against choosing repeated words following active responding to primes. In this present study, we explored conditions of prime processing that produce this pattern of results. Experiment 1 revealed that increased prime duration under passive viewing instructions produces the active priming pattern. Experiment 2 assessed memory for primes: With poor recognition of primes, there was a strong preference for repeated words; however, with good recognition of primes, this preference was eliminated. These results are modeled by a computational theory of optimal decision making, *responding optimally with unknown sources of evidence* (ROUSE). In ROUSE, a preference for repeated words results from source confusion between primes and choice words. A reversal in the direction of preference arises from the discounting of words known to have also appeared as primes.

Implicit priming is an ongoing area of interest in the study of human memory and perception. The term *priming* refers to the effect (usually a facilitation) of one stimulus (a *prime*) on a similar or identical stimulus (a *target*) presented at a later time. The term *implicit* refers to the nature of the task; facilitation is measured in tasks that can be accomplished by reference to stored knowledge and do not require explicit reference to primes. For example, participants might be asked to decide whether a letter string is a valid or invalid word (*lexical decision*). In another task, used in the present experiments, a word is briefly flashed and then masked, and participants attempt to identify the flashed target word (*perceptual identification*). In the latter task (and others), the prior presentation of identical or

related prime words can be made strategically irrelevant and nondiagnostic, and any effect of such presentations is therefore said to be implicit. The results of these implicit priming studies have proven highly useful in the development and testing of theories of memorial and perceptual processing.

Priming effects have been studied both with extended delays between presentation of prime and target (*long-term priming*) and with delays of only seconds or milliseconds (*short-term priming*). Both long-term and short-term priming have been shown to result in facilitation. Long-term priming typically produces effects only when the target is a repetition of a prime (e.g., Ratcliff, Hockley, & McKoon, 1985), whereas short-term priming has been observed to produce effects for a wide variety of similarity relations. Much of the short-term priming research has focused on associative/semantic relations (e.g., Evett & Humphreys, 1981; Marcel, 1983; Meyer & Schvaneveldt, 1971) and orthographic/phonemic similarity (e.g., Evett & Humphreys, 1981; Meyer, Schvaneveldt, & Ruddy, 1974). In addition, short-term orthographic/phonemic priming (e.g., Lukatela & Turvey, 1996) and repetition priming (e.g., Humphreys, Besner, & Quinlan, 1988) occasionally result in performance deficits. A recent study by

This research was supported by PHS-NIMH Grant 12717, NSF Grant SBR 9512089, NIMH Training Grant MH19879, NSF Grant KDI/LIS IBN-9873492, and ONR Grant N00014-00-1-0246. K.B.L. is now at the Department of Psychology, Yale University, New Haven, CT. Correspondence should be addressed to D. E. Huber, Department of Psychology, University of Colorado, 345 UCB, Boulder, CO 80309-0345 (e-mail: dhuber@psych.colorado.edu).

—Accepted by previous editorial team

Huber, Shiffrin, Lyle, and Ruys (2001) shed a good deal of light on the changeable nature of short-term priming. The present study followed closely on this work.

It is important to ask whether priming facilitation (or the opposite) results from changes in perceptual processing or from bias. For instance, in perceptual identification, improved accuracy in naming the flashed target may be due to improved acquisition of perceptual information from the target flash or a tendency to respond with recently seen items (i.e., a bias). This question has been looked at both with long-term priming (Bowers, 1999; Ratcliff & McKoon, 1997; Ratcliff, McKoon, & Verwoerd, 1989; Wagenmakers, Zeelenberg, & Raaijmakers, 2000; Zeelenberg, Wagenmakers, & Raaijmakers, 2002) and with short-term priming (Hochhaus & Johnston, 1996; Huber et al., 2001; Masson & Borowsky, 1998). In both situations, the results suggest that although both factors may play a role, bias is by far the larger of the two. One method used to disentangle bias from perceptual effects uses a two-alternative forced-choice (2AFC) version of perceptual identification (first introduced by Ratcliff et al., 1989). Following the flash of the target item, participants must choose between the target and a foil. In both long-term priming and short-term priming, exposure to the target facilitates performance, whereas prior exposure to the foil harms performance. This is an empirical measure of bias; there are both benefits and costs associated with priming.

Huber et al. (2001) used this paradigm to examine bias in short-term word priming (see Figure 3). There are four possible priming conditions in this paradigm: Neither choice word is primed (*neither-primed*), both choice words are primed (*both-primed*), only the target is primed (*target-primed*), and only the foil is primed (*foil-primed*). In order to test repetition priming with a both-primed condition, it is necessary to present two prime words. Huber et al. therefore used two primes in their experiments (and we did so in the present experiments as well). A bias to choose repeated or related words was observed when participants viewed prime words for 500 msec immediately prior to the target flash (i.e., $\text{target-primed} > \text{neither-primed} > \text{foil-primed}$). The authors termed this *passive* priming because participants were instructed that the prime words were a cue to prepare for the target flash but were not otherwise relevant for the task. Surprisingly, a slight bias *against* choosing repeated words (i.e., $\text{target-primed} < \text{neither-primed} < \text{foil-primed}$) was observed when participants were instructed to provide a response to prime words. This was termed *active* priming. Animacy, pleasantness, and part of speech were all used as active priming tasks with similar results. The flexible nature of the bias led the authors to adopt the term *preference* in place of bias, a choice that also helps avoid confusion with the term *bias* as it is used in signal detection theory. Another perplexing result was the observation of both-primed deficits ($\text{both-primed} < \text{neither-primed}$) occurring for both passive and active priming. Both the preference changes and the both-primed deficits were accounted for,

as explained next, by a computational theory based on Bayesian principles.

The theory, termed *responding optimally with unknown sources of evidence* (ROUSE), contains two mechanisms. The first is source confusion, consisting of confusion between features arising from primes and targets; this factor by itself produces a preference to choose words related to primes. In addition, the probabilistic nature of source confusion produces both-primed deficits (i.e., source confusion by chance sometimes favors one choice word and sometimes favor the other, adding noise to the decision process). The second mechanism is the attempt to overcome source confusion (on average) through the discounting of information that could have arisen from a prime instead of the target flash. This discounting of evidence allows a switch from a preference in favor of related choice words to a preference against related choice words, depending on how much discounting is employed by the participant. The preference changes result from the slight underestimation and overestimation of source confusion (an underestimate in passive priming and an overestimate in active priming).

In the present study, we followed up on this work by exploring the characteristics of primes and their processing that cause preference changes. In Experiment 1, we ascertained whether the passive/active distinction might result from different prime viewing durations. We tested this by varying prime duration in an experiment administered with passive priming instructions. The results indicated that preference does indeed change with prime viewing duration. Nevertheless, interpretation of this finding is problematic because participants are likely to engage in covert active processing of prime words given sufficiently long viewing durations. In Experiment 2, we assessed whether active priming instructions with brief, near-threshold, prime durations produce substantial discounting (i.e., a smaller than otherwise expected preference or even a reversal in preference). If so, the extent of discounting (i.e., the direction and magnitude of preference) may be a function of more than prime duration. In addition, in Experiment 2, we tested the assumption in ROUSE that knowledge of the primes (i.e., prime recognition) is necessary to guide the discounting process. Appropriate comparisons were achieved by breaking the data in half on the basis of high versus low prime recognition. Similar to the effect of longer prime viewing durations, it was revealed that high prime recognition produces a greater degree of discounting.

RESPONDING OPTIMALLY WITH UNKNOWN SOURCES OF EVIDENCE (ROUSE)

A Brief Summary

The description below is designed to present ROUSE in a simple and intuitive form. For a more detailed presentation, the reader is referred to Huber et al. (2001). ROUSE is a probabilistic model, and, in the original article, the

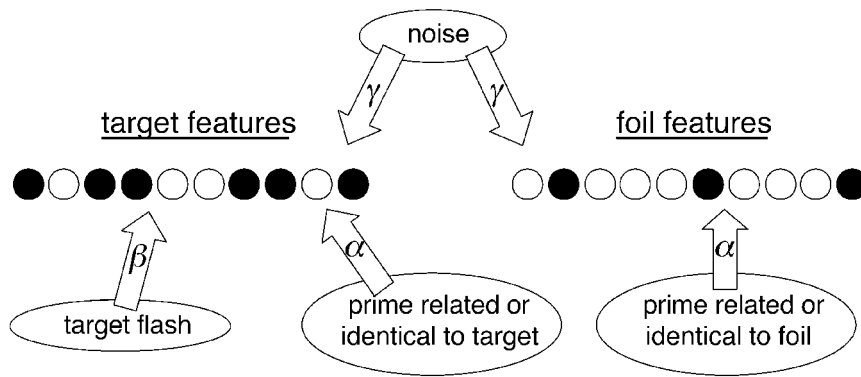


Figure 1. Source confusion through feature activation in ROUSE. Words are composed of features that exist in an active or an inactive state. In the perceptual identification paradigm, features may become active due to related or identical primes, the target flash, and/or noise (i.e., the visual mask), according to the corresponding activation parameters, α , β , and γ . Because it is unknown by which source a feature has become active, this situation results in source confusion. In particular, source confusion for prime-activated features results in a preference to choose prime-related words.

predictions were calculated through Monte Carlo simulations. The predictions in the present study were calculated analytically, greatly increasing the speed of the calculations (Huber, 2002).

Source confusion arises in ROUSE due to *unknown sources of evidence*. ROUSE supposes that words are composed of features and that the features of the choice words may become active, but the cause of the activation of at least some of the features might be ambiguous or unknown. Activation can arise due to presentation of related primes, the target flash, and/or noise. The probabilities α , β , and

γ correspond to prime, target, and noise activation (see Figure 1), and these values probabilistically determine which features of the target and the foil are active at the time of decision. Presenting related primes results in the activation of more features in a prime-related choice word and therefore produces a preference to choose prime-related words.

This preference is counteracted in ROUSE by the discounting of information, a factor that operates to produce “optimal responding.” Discounting in ROUSE occurs on a feature-by-feature basis. In making an optimal decision

		State of Activation	
		inactive (OFF)	active (ON)
Appeared in Prime(s)	NO	$\frac{(1-\gamma)(1-\beta)}{(1-\gamma)} = (1-\beta)$	$\frac{1 - (1-\gamma)(1-\beta)}{1 - (1-\gamma)}$
	YES	$\frac{(1-\gamma)(1-\alpha')(1-\beta)}{(1-\gamma)(1-\alpha')} = (1-\beta)$	$\frac{1 - (1-\gamma)(1-\alpha')(1-\beta)}{1 - (1-\gamma)(1-\alpha')}$

Figure 2. Discounting through feature likelihood ratios in ROUSE. These ratios are the probability that a feature exists in its on/off state given that it is part of the target versus the probability given that the feature is part of the foil. The important contingencies are for active (on) and inactive (off) features that did or did not appear in a prime (i.e., may or may not have been activated by a prime). The likelihood ratio for active features that appeared in a prime (i.e., the lower right-hand cell) represents a discounted level of evidence. As the estimate of prime activation, α' (i.e., discounting), increases toward 1.0, this ratio approaches the neutral value of 1.0 (and is less than the upper right-hand cell).

given the source confusion described above, a lower level of evidence (i.e., a *discounted* level of evidence) is given to features known to have been presented in prime words. Appropriate evidence levels are determined by calculating the likelihood that a feature is part of the target word, given that feature's state of activation and knowing whether that feature also appeared in a prime. The lower right cell in Figure 2 is the discounted likelihood ratio for an active feature that may have been activated by a prime; this ratio is closer to the neutral value of one, relative to the likelihood ratio in the upper right cell for an active feature that could not have been activated by a prime.

The likelihood ratio for a feature is the probability of the observed state of activation given that the target flash is a potential source of activation (i.e., that the feature is part of the target word) divided by the probability given that that the target flash is not a potential source of activation (i.e., that the feature is part of the foil word). The feature likelihood ratios in Figure 2 result from a 2×2 table of situations contingent on whether the feature is observed to be active (ON) or inactive (OFF) and whether or not a prime is a potential source of activation (i.e., whether the feature also appeared in a prime word). These feature likelihood ratios are then combined to provide two word likelihood ratios, one for each of the choice words, by multiplying the separate feature likelihood ratios of the features of a word. Finally, the word likelihood ratio for the target is divided by the word likelihood ratio for the foil to provide the odds in favor of the target. When the odds are above one, the target is chosen (a correct response); when the odds are less than one, the foil is chosen (an incorrect response). When the odds exactly equal one (i.e., the target and the foil are equally likely), a choice word is randomly selected.

As can be seen in Figure 2, the variable α' appears instead of α (the probability of feature activation due to a prime). This is because the system does not necessarily have access to the true probabilities of feature activation, and these probabilities must be estimated. The superscript ($\cdot$) indicates that an estimate of prime activation is used instead of the true value. Simulations have shown that misestimates of β and γ affect the model's performance in a quantitative, not qualitative, manner across the priming conditions; therefore, the estimates β' and γ' were set equal to β and γ . In contrast, the misestimate of prime activation (i.e., α' relative to α) is responsible for the magnitude of discounting and therefore the direction of preference. With passive priming, the system underestimates the probability of feature activation by primes (i.e., $\alpha' < \alpha$), resulting in a preference to choose prime-related words (essentially an underestimate of source confusion or too little discounting). With active priming, the system overestimates the probability of feature activation by primes (i.e., $\alpha' > \alpha$), resulting in a preference *against* prime-related words (essentially an overestimate of source confusion or too much discounting).

Regardless of the underestimation or overestimation of source confusion, ROUSE predicts that there will be performance deficits produced by priming (as seen most

clearly in the both-primed deficits and in the deficit in the average of the target-primed and foil-primed conditions, both relative to the neither-primed condition). These deficits are due to trial-by-trial variability in the number of features activated by primes. For instance, in the both-primed condition, the number of prime-activated features in either choice word will be the same, on average. However, on any particular trial, there may be more prime-activated features in one choice word than the other, due to the probabilistic nature of feature activation. This extra variability harms performance.

The Application of ROUSE to Experiments 1 and 2

Despite the success of ROUSE in accounting for the complex set of results found by Huber et al. (2001), the theory does not fully specify how several important empirical manipulations will affect the mechanisms of source confusion and discounting. In some aspects, the results reported here are exploratory and will help guide future elaborations of the theory to these manipulations. In other aspects, the results tightly constrain the theory in its present form and could potentially falsify the ROUSE model (e.g., see Roberts & Pashler, 2000, for a recent discussion on falsifying computational models). We will do our best to specify when the model is used to "describe" versus "predict" the results. We also note that a competitor model has been proposed by Ratcliff and McKoon (2001). Their multinomial model likewise contains mechanisms of source confusion and discounting and is expected to handle these new results about as well as the ROUSE model. It was not a goal of this study to differentiate between these models; for the most part, neither theory makes concrete predictions about the magnitude of source confusion and discounting for these particular experiments. The interested reader is referred to Huber, Shiffrin, Lyle, and Quach (in press) for experiments that highlight the differences between the two models.

The main reason for presenting the ROUSE model is to formalize, through the language of mathematics, the mechanisms of source confusion and discounting. In addition, applying ROUSE to these data allows us to examine the best-fitting parameters, which in several instances affords insight into the nature of the results that is not readily apparent from the behavioral data. An apt analogy comes from the application of signal detection theory (e.g., Macmillan & Creelman, 1991) to hit and false alarm data, in order to determine whether there has been a change in sensitivity, bias, or both. In the present case, we applied ROUSE in order to determine whether there had been a change in source confusion (α), discounting (α' relative to α), or both.

EXPERIMENT 1 Prime Duration

In Huber et al.'s (2001) experiments, it was not clear why passive priming resulted in an underestimation of source confusion and why active priming resulted in an overesti-

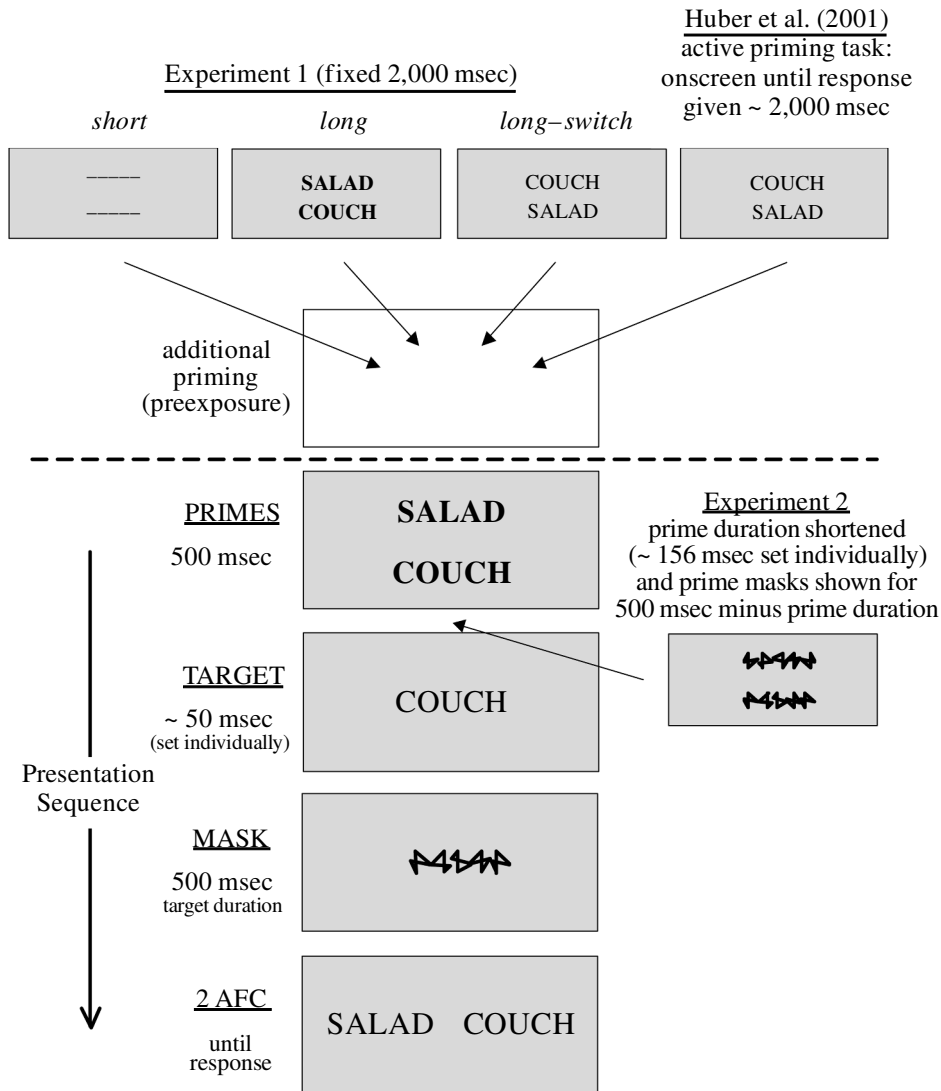


Figure 3. The sequence of events on a trial for the different experiments. This figure shows a both-primed trial with repetition priming. Besides this condition, experiments included neither-primed (e.g., primes TODAY and PRINT), target-primed (e.g., primes COUCH and PRINT), and foil-primed (e.g., primes TODAY and SALAD) conditions. In all experiments, target flash durations were determined individually in an attempt to set performance at 75 %. In the active priming condition of Huber et al. (2001), participants provided a response in relation to the primes. Following this, the primes switched locations and were displayed in boldface font for an additional 500 msec prior to the target flash. Only the second boldface 500-msec presentation of the primes was viewed in the passive priming condition. In Experiment 1, the average durations and types of presentations mimicked those found in Huber et al.'s (2001) passive priming condition (the short condition) and active priming condition (the long-switch condition) but without an active priming task. Experiment 2 was an active priming experiment with short, near-threshold presentations of the primes. In Experiment 2, prime presentation duration was set individually with an average time of 156 msec, and the primes were masked prior to the target flash. Following 2AFC responses in Experiment 2, the participants assessed whether each choice word also appeared as a prime word (i.e., prime recognition).

mation of source confusion. One hypothesis is based on the fact that the passive/active manipulation was confounded with prime duration (see Figure 3). In passive priming, primes were presented for 500 msec, whereas in

active priming, primes were presented until participants responded (taking at least 2,000 msec, on average). Following a response in the active priming task, the primes switched locations and were presented for an additional

500 msec in boldface prior to the target flash (such that the final 500 msec prior to the target flash were the same in both conditions).

In Experiment 1, we ascertained whether longer prime durations produce the active patterns of results in the absence of an active priming task. Using only passive priming instructions, three different types of prime presentations were used (see Figure 3). The first condition (termed the *short* condition) corresponded to the Huber et al. (2001) form of passive priming. Following a 2,000-msec presentation of “----” in the positions of the primes, the prime words appeared in boldface for 500 msec. The second condition (termed the *long* condition) tested the effect of prime duration, and primes were presented in boldface for 2,500 msec without a switch of positions. The final condition (termed the *long-switch* condition) mimicked the active priming condition used by Huber et al. (2001); following 2,000 msec of prime presentation in lightface, the primes switched locations and were presented in boldface for an additional 500 msec. Because Huber et al. (2001) found the largest changes in preference when the primes were identical to the choice words, repetition priming was used in Experiments 1 and 2.

Method

Participants. Thirty-three Indiana University undergraduates participated, receiving introductory psychology course credit for their participation. All participants were native English speakers, with normal or corrected-to-normal vision.

Materials. For all presentations, 1,000 five-letter words were used. These words had a minimum written language frequency of 4, as defined and measured by Kučera and Francis (1967). Randomly generated letter-like pattern masks were used to avoid pattern mask habituation (see Figure 3, for examples of the pattern masks). All words were displayed in uppercase Times Roman 22-point font.

Equipment. Stimulus materials were displayed on PC monitors, with presentation times synchronized to the vertical refresh. The re-

fresh rate was 120 Hz, providing display increments of 8.33 msec. In order to avoid phosphor decay, stimuli were displayed as black against a gray background. The participants' booths were enclosed, and the lighting was dim to avoid eyestrain. The resulting visual contrast was close to 100%. Monitor distance and font size were chosen such that both prime words and the center target word encompassed less than 3° of visual angle. All responses were collected through response boxes with four keys.

Procedure. All variables were within subjects. The basic design used the following two variables: priming condition, with four levels (neither-primed, both-primed, target-primed, and foil-primed), and type of prime presentation, with three levels (short, long, and long-switch). In the neither-primed conditions, two primes, a target, and a foil word were randomly selected from the pool of 1,000 five-letter words. Word selection occurred without replacement such that a given word appeared only on one trial within the experiment, thus avoiding contamination from long-term repetition priming. In the both-primed conditions, only two words were selected, because the primes were repeated as the choice words. For the target-primed and foil-primed conditions, the target or foil was accordingly repeated, and the other choice word was randomly selected (these conditions required three words).

As can be seen in Figure 3, two prime words were presented on every trial: one above and one below the fixation point. The short, long, and long-switch prime presentations were as described in the introduction to Experiment 1 (and are shown in Figure 3). During the 2AFC, the target and the foil were presented on the left and the right of fixation. The top/down position of the primes and the left/right position of the choice words were both fully counterbalanced. On every trial, a sequence of events occurred as shown in Figure 3. Prior to the first display of the prime words (i.e., the first screen in Figure 3), a fixation point was displayed for 500 msec, followed by a blank screen for 500 msec.

The experiment began with 16 trials of practice using short prime presentations, and priming was always neither-primed. The participants were told that prime words were a warning to prepare for the flash of the target word (i.e., passive priming instructions). Feedback was given on every trial throughout the entire experiment. Target flash duration was set at 100 msec during practice, allowing near-perfect performance. Following practice, the participants received a

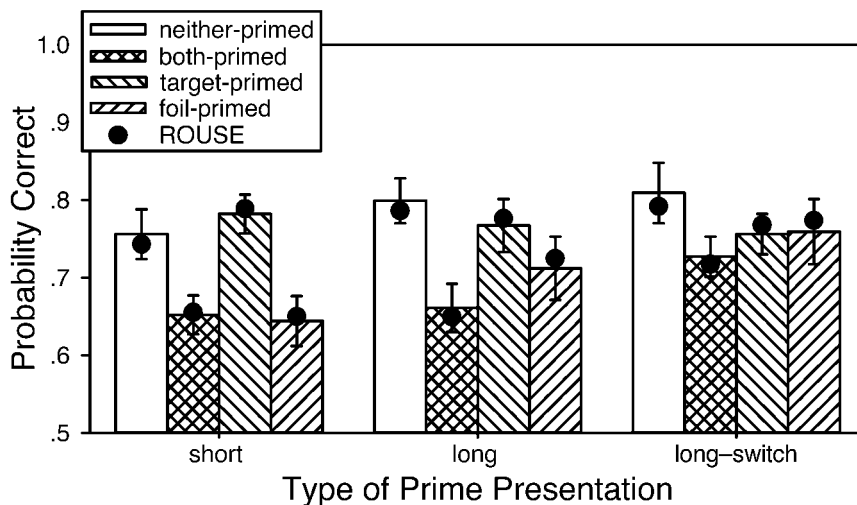


Figure 4. Observed and predicted results for Experiment 1, a manipulation of prime duration. The error bars are 2 standard errors of the mean. The ROUSE predictions are the result of the best-fitting parameters appearing in Table 1.

block of 64 trials with neither-primed short prime presentations. The purpose of these trials was to find the target flash duration at which performance was approximately 75%. Resultant durations averaged 43 msec, although there were large individual differences, with times ranging from 33 to 117 msec. A staircase method was used to find the appropriate target flash duration. The participants were fully informed about this procedure. Following these 64 trials, two blocks of 96 experimental trials were presented. At the start of the first experimental block, the participants were informed that the warning words (i.e., the primes) would often repeat as the target and foil words. Furthermore, they were informed that the target word would be a repeat just as often as the foil word and that there could be no effective strategy in terms of choosing or not choosing repeated words. The entire experiment took approximately 40 min.

Results

The results are depicted in Figure 4, and the means and standard errors can be found in the Appendix. There was a main effect of the four priming conditions [$F(3,96) = 9.57$, $MS_e = 0.258$, $p < .001$] and a main effect of the three types of prime presentations [$F(2,64) = 4.62$, $MS_e = 0.0977$, $p < .025$]. These two variables interacted [$F(6,192) = 2.96$, $MS_e = 0.034$, $p < .01$]. Before discussing the nature of this interaction, we will present statistical techniques for analyzing the data (these techniques are discussed in greater detail in Huber et al., 2001).

We assessed decisional effects (i.e., the existence of preference) by comparing conditions in which the foil was primed with conditions in which the foil was not primed, while holding target priming constant (i.e., foil-primed compared with neither-primed, and both-primed compared with target-primed). Such comparisons hold constant any perceptual effects of priming the target (should they exist) and assess preferential effects of priming the foil (by *preferential* we mean any effect on performance that is independent of the target flash and therefore applies equally to all prime words). We combined the two comparisons into a single F test to assess the existence of preferential effects. If a preferential effect was found, we then assessed the direction of preference with a t test comparing the target-primed and foil-primed conditions.

This test for the existence of preferential effects is not only sensitive to a directional effect of preference (i.e., a preference for or against prime-related words) but is also sensitive to changes in performance as might be induced by preferential variability (i.e., variability in the extent of preference across trials and words). This variability reduces performance in the both-primed condition and reduces the average of the target-primed and foil-primed conditions. Therefore, in some situations, the test for the existence of preferential effects might suggest that preference played a role, and the test for the direction of preference might show that the preference was neutral in its direction (i.e., on some trials, the preference is for prime-related words, and, on other trials, the preference is against prime-related words, and these average out across trials but reduce performance). Evidence of preferential variability is more directly tested through a comparison of the both-primed condition with the neither-primed condi-

tion (depending on one's viewpoint, this test could also be considered a test of inhibition for repeated words).

Moving from left to right in Figure 4, with short prime presentations, there was a preferential effect [$F(1,32) = 17.17$, $MS_e = 0.485$, $p < .001$], and the nature of the preference was to choose repeated words [$t(32) = 3.72$, $SE = 0.0372$, $p < .001$]. In addition, there was a both-primed deficit [$t(32) = 2.97$, $SE = 0.0351$, $p < .005$], suggesting preferential variability was a factor. With long prime presentations, there was a preferential effect [$F(1,32) = 15.67$, $MS_e = 0.308$, $p < .001$], although the direction of preference was neutral [$t(32) = 1.25$, $SE = 0.0438$, $p = .22$]. Presumably, the preferential effect was due to variability, as demonstrated by the both-primed deficit [$t(32) = 4.55$, $SE = 0.0304$, $p < .001$]. With long-switch prime presentations, there was a preferential effect [$F(1,32) = 4.29$, $MS_e = 0.05$, $p < .05$] that, on average, was neutral [$t(32) = 0.1$, $SE = 0.0386$, $p = .92$]. As with the other types of prime presentations, there was a both-primed deficit [$t(32) = 3.06$, $SE = 0.0266$, $p < .005$].

Discussion

Experiment 1 demonstrates that the extent of discounting (i.e., the extent of preference removal) is correlated with prime duration; in the long condition and especially in the long-switch condition, there was a neutral preference, in contrast to the robust preference for repeated words in the short condition. In light of this preference reduction with increasing prime duration (or also thought of as increased discounting with increased prime duration), it is not necessary to use an active priming task in order to induce the active priming pattern of results. Although longer prime durations failed to reverse the direction of preference, Huber et al. (2001) did not observe preference reversals on all occasions with active repetition priming (in one experiment, the preference was neutral, and, in another experiment, the preference reversal fell just shy of significance). Therefore, it might be that prime duration is solely responsible for most of the observed effects. However, it is likely that participants use the extended prime-viewing time to engage in more active processing of the primes. Additional experiments are necessary to disentangle the time hypothesis from the processing hypothesis.

Regardless, the present results shed some light on the role of prime duration in experiments using more traditional priming measures. Traditional priming measures, such as lexical decision, speed of naming, and the more traditional naming form of perceptual identification, are most closely akin to a comparison of the target-primed and neither-primed conditions in our 2AFC paradigm (i.e., traditional measures compare a primed condition with an unprimed condition and do not measure the decisional effect of primed foil items). As can be seen in Figure 4, the target-primed condition is above the neither-primed condition for the short prime duration, whereas it is below the neither-primed conditions for the two types of long prime duration. As such, with repetition priming, it might be expected that

increased prime duration results in decreased prime facilitation or even priming deficits.

Humphreys et al. (1988) examined repetition priming as measured with naming accuracy in a perceptual identification task (i.e., participants spoke aloud the names of the briefly flashed words). Previously, Evett and Humphreys (1981) had demonstrated a priming facilitation for repetition priming in this paradigm with short (~ 40-msec) sub-threshold prime presentations. Humphreys et al. replicated the repetition priming advantage with brief prime presentations and extended the set of conditions to include longer (300 msec) prime presentations. Surprisingly, they found a large negative effect for repetition priming at the longer duration. In a subsequent experiment, they were able to recover the positive effect of repetition priming by breaking the 300 msec between prime onset and target onset into a 180-msec prime presentation (still well above threshold) followed by a 120-msec mask presentation. They interpreted this change in terms of additional cues for distinguishing the prime and the target as separate visual events. Our results suggest an alternative interpretation: 300-msec priming led to overdiscounting and excessive removal of target preference, but 180 msec followed by masking led to underdiscounting and inadequate removal of target preference.

ROUSE Predictions and Discussion

Fitting ROUSE to Experiment 1 provides an opportunity to gain additional insight into the empirical results by quantifying the levels of source confusion and discounting.

Because each of the three types of prime presentation are likely to induce different levels of target perception (β), source confusion (α), and discounting (α relative to α'), these three parameters were allowed different values for each of the types of prime presentation. As in Huber et al. (2001), the probability of noise activation (γ) was fixed at the low level of .02, and the number of features contained in each word was fixed at 20. With no parameter applying across the three types of prime presentations, the fitting routine was run separately on each type, and the separate error scores are reported in Table 1.

The results obtained using the best-fitting parameters are seen as the dots in Figure 4, and the exact values for each

condition can be found in the Appendix. The best-fitting parameters are given in Table 1. Upper and lower confidence limits were found for each parameter (see Huber, 2002, for a generally applicable technique for producing parameter confidence limits), with the chosen level of confidence based on $\Delta\chi^2 = 16.0$, which corresponds to the "true" parameter value lying between the confidence limits with a probability in excess of .9999.

As can be seen in Figure 4, the fit is very good. A key prediction tested by fitting ROUSE is that the average of the target-primed and foil-primed conditions will be roughly equal to the average of the neither-primed and both-primed conditions, for a given type of prime presentation. Source confusion causes half as much interference in the target-primed and foil-primed conditions, as compared with the both-primed condition, because only one of the two choice words has been primed, instead of both. Therefore, for close to accurate estimates of source confusion (α'), the model predicted this relationship, and the data consistently conformed to the prediction. If the observed results had been otherwise (e.g., the average of the target-primed and foil-primed conditions above that of the neither-primed condition), the model would have been falsified.

Examining the best-fitting parameters in Table 1 provides further insight into the prime duration manipulation. Given the range of acceptable α values indicated by the confidence limits, there is no clear trend for source confusion across the three types of prime presentation. However, there is a clear trend in discounting (i.e., the relationship between α and α'). There is too little discounting for the short condition (the usual result with passive priming with an underestimation of α), close to optimal discounting (accurate estimation of α) in the long condition, and excessive discounting (similar to the usual active priming result with an overestimation of α) in the long-switch condition. Beyond source confusion and discounting, an examination of the best-fitting parameters also reveals that the three types of prime presentation resulted in different degrees of target perceptibility (β). For reasons perhaps relating to attentional capture or differential forward masking, long duration primes interfered less with the ability to perceive the target flash.

Table 1
Best-Fitting ROUSE Parameters With Lower and Upper 99.99% Confidence Limits

Duration/Recognition	Experiment 1			Experiment 2	
	Short	Long	Long-Switch	Low	High
α (prime actual)	.052 .074 .103	.080 .144 .245	.029 .049 .081	.022 .031 .043	
α' (prime estimate)	.046 .059 .074	.082 .130 .201	.048 .075 .115	.004 .006 .011	.020 .029 .043
β (target flash)	.041 .045 .049	.054 .057 .060	.055 .059 .063	.063 .065 .068	
$\Sigma\chi^2$ (error)	.728	1.510	2.247	2.070	

Note—Parameters in boldface are the best-fitting values, with lower and upper parameter confidence limits appearing to the left and right in lightface. In Experiment 2, the same α and β parameter were used for both low and high prime recognition. Therefore, these parameters are listed between the low and high columns. The chi-squared error ($\Sigma\chi^2$) was calculated using the best-fitting parameter values. In Experiment 2, the fitting routines were simultaneously applied to low and high prime recognition, resulting in a single $\Sigma\chi^2$ across these conditions.

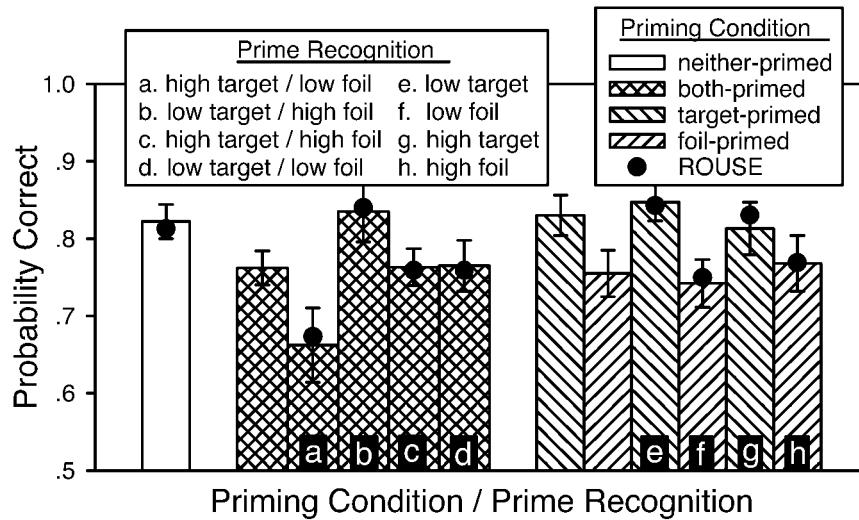


Figure 5. Observed and predicted results for Experiment 2, including a breakdown of the data based on high and low prime recognition of repeated words (see Experiment 2 Results section for a description of the breakdown). The error bars are 2 standard errors of the mean. The ROUSE predictions are the result of the best-fitting parameters appearing in Table 1.

EXPERIMENT 2 Prime Recognition

Experiment 1 demonstrated that prime duration and discounting are strongly correlated. However, longer prime presentations may naturally induce a higher, more active, degree of prime processing (e.g., given sufficient time, participants may dwell on the meaning of the primes). To help disentangle these possibilities, in Experiment 2, we ascertained whether an active priming task results in substantial discounting despite the use of short prime durations. In this experiment, we used brief, near-threshold, primes and included prime recognition as an active priming task.

In addition to testing for discounting with brief primes, we evaluated the relation between discounting and the ability to recognize the primes. ROUSE the discounting of features is guided by the knowledge of which features also appeared in prime words. If the participant does not recognize that a choice word is a repetition of a prime, then it seems less likely that the word identification system could know which features to discount. This would be revealed by a lowered estimate of discounting when primes are not recognized. Therefore, the active priming task was to judge whether each choice word also appeared as a prime. These judgments were given immediately following the usual 2AFC response concerning which word appeared during the target flash.

Although not a specific quantitative prediction of ROUSE, the failure to observe the expected relationship between discounting and prime recognition would cause us to question whether the ROUSE implementation of discounting was appropriate.

Method

Forty-two University of Colorado undergraduates received introductory psychology course credit for their participation. Unless otherwise noted, all materials, equipment, and procedures were the same as those used in Experiment 1. Computer monitor refresh rates were 70 Hz, providing display increments of 14.29 msec. The participants were tested in a single well-lit room with dividers between each setup. Responses were collected through computer keyboard.

The design of Experiment 2 was relatively simple. The only conditions were the four priming conditions: neither-primed, both-primed, target-primed, and foil-primed. The order of conditions across trials was determined randomly for each participant, and all conditions appeared equally often. Left/right choice word position and top/down prime position of the target were fully crossed variables. As in Experiment 1, repetition priming was used. The lettered conditions appearing in Figure 5 and Table 2 are the result of a post hoc breakdown of the data based on the prime recognition responses.

In order to facilitate a breakdown of the data based on high versus low prime recognition, both types of responses were collected on an 8-point rating scale. These were obtained through the use of an onscreen slider bar with four positions on each side of a neutral middle position. For the first response concerning choice word as target (i.e., 2AFC), the response scale appeared horizontally, and the par-

Table 2
Probability of Rating Choice Words as Primes
(Prime Recognition) in Experiment 2

Choice Word	Priming Condition	Hits	False Alarms	SE
Target	Neither	—	.260	.022
	Both	.784	—	.022
	Target	.776	—	.022
	Foil	—	.218	.027
Foil	Neither	—	.182	.020
	Both	.809	—	.022
	Target	—	.168	.020
	Foil	.742	—	.024

ticipants moved the position of the slider with left/right arrow keys toward one choice word or the other. A neutral response was not allowed, such that once the slider left its initial position in the middle, the slider could not return to the middle neutral position. If the participants moved the slider bar back to the middle position, it would jump to the opposite side. The participants were instructed to slide the bar toward the chosen word, with the number of steps corresponding to their certainty. If guessing, they were instructed to move the slider only a single step in the chosen direction.

Following this 2AFC decision (the final position being recorded with a press of the "enter" key), the central horizontal response scale disappeared and was replaced by a vertical response scale next to the left choice word. This next response was whether the left choice word appeared as a prime, with up being "yes" and down being "no" (the slider was moved with the up/down arrow keys). Because left/right position of the target word was a fully counterbalanced variable, always testing the left before the right choice word did not introduce any confounds. Similar to that for the first response, the slider bar had four positions above and four below the central position, which served only as the starting point. Again, confidence was the variable determining how far to move the slider. Following this response, a similar vertical response scale appeared next to the right choice word and the same choice-word-as-prime response given (both choice words remained onscreen throughout the recording of all three responses). After this last response, feedback appeared, with correctness listed for each of the three responses (correctness was determined on the basis of the midpoint of the response scales).

In order to keep prime recognition below ceiling, we used short prime durations. Furthermore, a screen containing masks in the same positions as the two prime words was displayed following presentation of the primes and just prior to presentation of the target flash (see Figure 3). The duration of these prime masks was 500 msec minus the duration of the primes. Similar to target durations, prime durations were set individually, such that performance on the prime recognition responses was close to 75%. On average, across participants, the appropriate prime duration was 156 msec, although there was tremendous variability, with times ranging from 29 to 343 msec. Target flash durations ranged from 14 to 186 msec, and the average was 61 msec.

As in Experiment 1, there were 16 practice trials, followed by a block of 64 trials in which flash durations were adjusted. Unlike in Experiment 1, both prime duration and target duration were independently varied during these 64 trials. Following this threshold determination block of trials (which were all neither-primed), the participants received two experimental blocks, with 48 trials in each block. Because three responses were needed on each trial, the total duration of the experiment was comparable to that of Experiment 1. It should be emphasized that, unlike the passive priming instructions of Experiment 1, the participants were instructed in this experiment to pay close attention to the briefly presented prime words because their memory for these words would be tested in the second and third responses given on each trial.

Results and Discussion

Prior to breaking down the 2AFC data based on the prime recognition responses, we report statistics on the collapsed data in regard to the existence and direction of preference (see the discussion at the start of the Experiment 1 Results section for a description of these statistics). The collapsed data appear as the nonlettered conditions in Figure 5 and the Appendix. There was a main effect of the four priming conditions [$F(3,123) = 4.17, MS_e = 0.0654, p < .01$]. Assessing the nature of this priming, there were preferential effects [$F(1,41) = 8.16, MS_e = 0.194, p < .01$]. Com-

paring the target-primed and foil-primed conditions [$t(41) = 1.87, SE = 0.0404, p = .07$, two-tailed], the apparent preference to choose repeated words (see Figure 5) did not reach significance. In addition, there was a both-primed deficit [$t(41) = 3.41, SE = 0.0178, p < .0025$].

Next, the prime recognition data was used to parse the 2AFC results. In the case of false alarms (i.e., labeling a nonprimed word as having been a prime), it was unclear how discounting should apply. It could be argued that the features of nonprimed choice words judged to be primes should be discounted. However, in most cases, such false alarms occurred with low confidence, a factor that might lead to little discounting. Given that false alarms occurred infrequently (as can be seen in Table 2) and that high-confidence false alarms were exceedingly rare, discounting effects for false alarms would probably be hard to verify. Indeed, use of the parsing technique described below, applied to nonprimed words, revealed little difference between the 2AFC data for high and low prime recognition of non-primed words. For this reason, the reported post hoc parsing of the 2AFC data based on prime recognition is restricted to primed choice words (i.e., hits and misses).

We partitioned the hits and misses on the basis of ratings in a way that would provide sufficient numbers to carry out analyses: Separate prime recognition criteria were determined for the ratings of primed words (collapsing across both hits and misses), such that half of the 24 trials for each condition per participant were labeled as having resulted in high prime recognition and the other half as having resulted in low prime recognition. These criteria were determined separately for each primed choice word (i.e., in the both-primed conditions, there were two such criteria: one for the target and one for the foil). If two or more trials fell on the criterion, it was randomly determined whether these trials were high or low prime recognition, so as to maintain 50% of the data in each category. The 2AFC data were then parsed according to level of prime recognition, resulting in a total of eight new post hoc conditions (letters *a-h* in Figure 5 and the Appendix). In the both-primed condition, both the target and the foil could be separately rated as high or low prime recognition, giving rise to four post hoc conditions labeled *a-d*. For the target-primed condition, only the target choice could be parsed by prime recognition, giving rise to the post hoc conditions *e* and *g*. Similarly, the foil-primed condition was parsed into post hoc conditions *f* and *h*. This method of breaking down the data does not guarantee that the same number of trials will fall into each of the four both-primed post hoc conditions (*a-d*), although it does guarantee equal numbers for the target-primed and foil-primed post hoc conditions (*e-h*). The actual numbers of trials that fell into each post hoc condition are found in the Appendix.

Across the four post hoc both-primed conditions (*a-d*), there were significant differences [$F(3,120) = 5.21, MS_e = 0.207, p < .0025$]. These occurred as an effect of high/low prime recognition for the target [$F(1,40) = 7.64$,

$MS_e = 0.313, p < .01$], as well as an effect of high/low prime recognition for the foil [$F(1,40) = 6.51, MS_e = 0.303, p < .025$]. These two effects did not interact [$F(1,40) = 0.17, MS_e = 0.0054, p = .68$]. As can be seen in Figure 5, high prime recognition for the target resulted in worse performance (i.e., compare *a* with *d* and compare *c* with *b*), whereas high prime recognition for the foil resulted in better performance (i.e., compare *b* with *d* and compare *c* with *a*), and both of these effects did not depend on the level of prime recognition for the other choice word (i.e., no interaction). This pattern is to be expected if high prime recognition results in greater discounting of the recognized word. In addition, the lack of an interaction strongly suggests that discounting is separately, and independently, applied to each primed choice word.

Next, the four single-primed post hoc conditions (*e-h*) are considered. Although direct comparisons of high/low prime recognition for the target-primed condition (*e* vs. *g*) [$t(41) = 1.22, SE = 0.0276, p = .22$] and for the foil-primed condition (*f* vs. *h*) [$t(41) = 0.91, SE = 0.0285, p = .37$] revealed no effect of prime recognition, comparing the target-primed conditions with the foil-primed conditions was more informative. When recognition for the primed word was low, there was a preference to choose repeated words [$t(41) = 2.71, SE = 0.0389, p < .01$], revealed by comparing the target-primed and foil-primed conditions (*e* vs. *f*). In contrast, when recognition for the primed word was high (*g* vs. *h*), the preference was neutral in its direction [$t(41) = 0.89, SE = 0.0511, p = .38$]. Essentially, the nearly significant preference for repeated words in the collapsed data breaks down as a strong preference with low prime recognition and a neutral preference with high prime recognition.

The main finding of Experiment 2 is that the ability to recognize a choice word as having been a prime word (i.e., prime recognition) was directly related to how strongly that choice word was mistrusted (i.e., the extent of discounting). Similar to increases in prime duration (Experiment 1), increases in prime recognition corresponded with greater discounting and therefore diminished preference for repeated words. The finding of a direct relationship between prime recognition and discounting suggests that prime viewing duration is not the sole variable determining the extent of discounting. This is further evidenced in the combined data: Collapsing across prime recognition level, the preference was revealed to be neutral in its direction despite the use of short prime durations.

ROUSE Predictions and Discussion

The results of the best-fitting parameters in Table 1 appear as the ROUSE predictions in Figure 5. The predictions are also listed numerically in the Appendix. Simulations were run in the same manner as discussed in Experiment 1. Collapsed data were not directly fit, considering that accurate fits to the post hoc conditions can be recombined to produce the collapsed data. Despite using only four free parameters, the fit to the data is remarkably good. The difference between high and low prime recog-

nition was captured through the use of one estimate of source confusion (α') for high prime recognition and a second for low prime recognition. Regardless of the level of prime recognition, the same levels of source confusion (α) and target perception (β) were used. As might be expected, the best-fitting parameter for discounting for high prime recognition was larger than for low prime recognition.

It is noteworthy that the model fits the data well, assuming that only discounting varies with changes in prime recognition. In ROUSE, source confusion for prime features (α) is a separate process from discounting of evidence for features known to have been in primes (α'). We fit the prime recognition breakdown by assuming a fixed level of α and α' variable level of α' . However, it might be reasonable to assume that recognizable primes are more readily parsed as separate visual events from the flash, and, therefore, source confusion should be less in these cases. If this were the situation, different levels of α would produce different performance levels for the both-primed conditions in which recognizability varied (labeled *c* and *d* in Figure 5). Instead, these conditions produced identical performance. Combined with the good fit of the present version of the model, this result suggests that prime recognizability affects only discounting (α'), but not source confusion (α).

It should be noted that a different, but equally accurate, account of these data could be provided by the multinomial model of Ratcliff and McKoon (2001). In that model, both source confusion and discounting contribute to the magnitude of both-primed deficits. A fit to these data with the multinomial model would necessarily assume that higher prime recognition caused both increased discounting and decreased source confusion. In this manner, conditions *c* and *d* could be equated due to a tradeoff between source confusion and discounting. Such an explanation is not available within ROUSE because the both-primed deficit is a function only of source confusion. Again, we refer the interested reader to Huber et al. (in press) for experiments that more directly compare and contrast the two models.

GENERAL DISCUSSION

The reported experiments help determine important variables for producing change in the preference to choose repeated (and, by extension, related) words in a perceptual identification task. The ROUSE model assumes that these changes result from differential degrees of discounting for primed words. Two different effects of prime viewing were observed. In Experiment 1, it was found that increased prime exposure duration (which also increased the prime-target SOA) resulted in diminished preference for repeated words (i.e., greater discounting). In Experiment 2, with near-threshold prime durations, it was found that good prime recognition was associated with diminished preference for repeated words. Because Experiment 2 provided evidence that prime recognition is an important component of discounting, it is unclear whether the increased prime du-

ration in Experiment 1 directly caused increased discounting or whether increased prime recognition was a mediating factor. Moreover, an active priming task, such as that employed by Huber et al. (2001), can be expected to increase prime recognition beyond any increases due to prime duration and therefore further decrease (and perhaps reverse) preference, relative to passive prime viewing.

These results were shown to be in good quantitative and qualitative agreement with the ROUSE model. ROUSE predicted, and the data verified, that the average of the neither-primed and both-primed conditions should be roughly equal to the average of the target-primed and foil-primed conditions (equivalently, the both-primed condition should suffer twice as much harm due to priming, as compared with the target-primed and foil-primed conditions). Although ROUSE made no clear prediction regarding the role of prime duration, the finding in Experiment 1 that long prime durations produce more discounting will help constrain future elaborations of the theory. In Experiment 2, we tested and confirmed a stronger prediction of ROUSE that greater prime recognition would result in greater discounting. This prediction derives from the underlying assumption in ROUSE that discounting is applied only to features known to have been contained in primes.

Huber et al. (2001) stressed the significance of preference reversals in explaining many seemingly contradictory results found in the short-term priming literature. Because the direction of preference changes across conditions, it is important to include foil-primed conditions in priming studies; these allow assessment of the direction and magnitude of preference independent of perceptual effects. The present study follows up on this work and begins the process of specifying key variables that result in preference changes.

REFERENCES

- BOWERS, J. S. (1999). Priming is not all bias: Commentary on Ratcliff and McKoon (1997). *Psychological Review*, **106**, 582-596.
- EVETT, L. J., & HUMPHREYS, G. W. (1981). The use of abstract graphemic information in lexical access. *Quarterly Journal of Experimental Psychology*, **33A**, 325-350.
- HOCHHAUS, L., & JOHNSTON, J. C. (1996). Perceptual repetition blindness effects. *Journal of Experimental Psychology: Human Perception & Performance*, **22**, 355-360.
- HUBER, D. E. (2002). *Computer simulations of the ROUSE model: An analytic method and a generally applicable technique for producing parameter confidence intervals*. Manuscript submitted for publication.
- HUBER, D. E., SHIFFRIN, R. M., LYLE, K. B., & QUACH, R. (in press). Mechanisms of source confusion and discounting in short-term priming: 2. Effects of prime similarity and target duration. *Journal of Experimental Psychology: Learning, Memory, & Cognition*.
- HUBER, D. E., SHIFFRIN, R. M., LYLE, K. B., & RUYS, K. I. (2001). Perception and preference in short-term word priming. *Psychological Review*, **108**, 149-182.
- HUMPHREYS, G. W., BESNER, D., & QUINLAN, P. T. (1988). Even perception and the word repetition effect. *Journal of Experimental Psychology: General*, **117**, 51-67.
- KUČERA, H., & FRANCIS, W. (1967). *Computational analysis of present day American English*. Providence, RI: Brown University Press.
- LUKATELA, G., & TURVEY, M. T. (1996). Inhibition of naming by rhyming primes. *Perception & Psychophysics*, **58**, 823-835.
- MACMILLAN, N. A., & CREELMAN, C. D. (1991). *Detection theory: A user's guide*. New York: Cambridge University Press.
- MARCEL, A. J. (1983). Conscious and unconscious perception: Experiments on visual masking and word recognition. *Cognitive Psychology*, **15**, 197-237.
- MASSON, M. E. J., & BOROWSKY, R. (1998). More than meets the eye: Context effects in word identification. *Memory & Cognition*, **26**, 1245-1269.
- MEYER, D. E., & SCHVANEVELDT, R. W. (1971). Facilitation in recognizing pairs of words: Evidence of a dependence between retrieval operations. *Journal of Experimental Psychology*, **90**, 227-234.
- MEYER, D. E., SCHVANEVELDT, R. W., & RUDDY, M. G. (1974). Functions of graphemic and phonemic codes in visual word-recognition. *Memory & Cognition*, **2**, 309-321.
- RATCLIFF, R., HOCKLEY, W., & MCKOON, G. (1985). Components of activation: Repetition and priming effects in lexical decision and recognition. *Journal of Experimental Psychology: General*, **114**, 435-450.
- RATCLIFF, R., & MCKOON, G. (1997). A counter model for implicit priming in perceptual word identification. *Psychological Review*, **104**, 319-343.
- RATCLIFF, R., & MCKOON, G. (2001). A multinomial model for short-term priming in word identification. *Psychological Review*, **108**, 835-846.
- RATCLIFF, R., MCKOON, G., & VERWOERD, M. (1989). A bias interpretation of facilitation in perceptual identification. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, **15**, 378-387.
- ROBERTS, S., & PASHLER, H. (2000). How persuasive is a good fit? A comment on theory testing. *Psychological Review*, **107**, 358-367.
- WAGENMAKERS, E.-J. M., ZEELLENBERG, R., & RAAIJMAKERS, J. G. W. (2000). Testing the counter model for perceptual identification: Effects of repetition priming and word frequency. *Psychonomic Bulletin & Review*, **7**, 662-667.
- ZEELLENBERG, R., WAGENMAKERS, E.-J. M., & RAAIJMAKERS, J. G. W. (2002). Priming in implicit memory tasks: Prior study causes enhanced discriminability, not only bias. *Journal of Experimental Psychology: General*, **131**, 38-47.

APPENDIX
Observed (Obs) and Predicted (Pred) 2AFC Data

Prime Duration/ Prime Recognition	Priming Condition	Data		SE	N
		Obs	Pred		
Experiment 1					
Short	Neither	.756	.743	.032	528
	Both	.652	.655	.025	528
	Target	.782	.789	.025	528
	Foil	.644	.650	.032	528
Long	Neither	.799	.786	.029	528
	Both	.661	.650	.031	528
	Target	.767	.776	.034	528
	Foil	.712	.725	.041	528
Long-Switch	Neither	.809	.792	.039	528
	Both	.727	.717	.026	528
	Target	.756	.768	.026	528
	Foil	.759	.774	.042	528
Experiment 2					
Collapsed	Neither	.822	.813	.022	1,008
Collapsed	Both	.762		.022	1,008
a. High target/low foil	Both	.662	.673	.048	175
b. Low target/high foil	Both	.835	.840	.039	175
c. High target/high foil	Both	.763	.759	.024	329
d. Low target/low foil	Both	.765	.759	.033	329
Collapsed	Target	.830		.026	1,008
Collapsed	Foil	.755		.030	1,008
e. Low target	Target	.847	.843	.024	504
f. Low foil	Foil	.742	.750	.031	504
g. High target	Target	.813	.830	.034	504
h. High foil	Foil	.768	.769	.036	504

(Manuscript received November 28, 2000;
revision accepted for publication February 12, 2002.)