

Perception and Preference in Short-term Word Priming

David Ernest Huber

Submitted to the faculty of the University Graduate School
in partial fulfillment of the requirements
for the degree
Doctor of Philosophy
in the Department of Psychology
Indiana University

January 2000

c 2000

David E. Huber

ALL RIGHTS RESERVED

Acknowledgements

Many people were instrumental to the development, execution, and writing of this thesis. Most important among these is my advisor, Richard Shiffrin. His insight and guidance are sprinkled liberally throughout this work. Nearly all the experiments were designed and performed through the hard work and dedication of two students: Keith Lyle and Kirsten Ruys. Their names, as well as Richard Shiffrin's, appear in a version of this work submitted for publication. Besides those directly involved in the priming studies, there was a supporting cast of characters. Most important amongst these were Karen Loffland and Coralee Sons who made all the red tape melt away. Without the expertise of Bill Wang the technical difficulties would have been insurmountable. The advice and comments of members of the Shiffrin lab, which included Mark Steyvers, David Diller, Denis Cousineau, Amy Criss, Lael Schooler, Rachel Shoup and Peter Nobel, were invaluable. William K. Estes and Eric-Jan Wagenmakers deserve recognition for helpful comments on early drafts of this thesis. I am grateful to my parents, Ernest and Ellen, for bringing me into this world and instilling in me a sense of scientific curiosity. Most important of all, I would like to thank my wife, Christina Anderson, who believed in me through it all.

Preface

This thesis presents a new theory of short-term priming termed ROUSE, standing for Responding Optimally with Unknown Sources of Evidence. Short-term priming refers to paradigms in which ‘irrelevant’ prime(s) are presented immediately prior to a target presentation to which a response must be given; typically the task requires a lexical decision or naming response (measured by response time, used when the target is above threshold), or identification (measured by accuracy, when the target is presented at threshold). Associative, orthographic/phonemic, and repetition priming are considered. The new theory is closely tied to the results from a new set of studies that considerably expand the set of conditions tested in such paradigms. We believe the results would appear inexplicable without the associated theory. Conversely, the theory would be hard to justify without reference to the results. These considerations lead us to delay presentation of the theory until the results of the first study are presented.

A central theme of the thesis is the attempt to understand the effect of a prime upon performance. In particular, we are interested in distinguishing effects that alter the perceptual response to the target during and shortly after its presentation from preference effects that alter other aspects of the priming situation. These are subtle distinctions (e.g. both perceptual and preference effects can affect bias and sensitivity in signal detection terms); their understanding requires a review of empirical and theoretical research, as well as detailed analysis of our present results. Such considerations have led us to organize the thesis in the following way. The introduction reviews the most pertinent prior empirical findings and theoretical interpretations, and relates our notions of perception and preference to the notions of sensitivity and bias that are found in signal detection theory. The first study is then presented; its results are used to motivate the ROUSE theory, presented next. The remaining studies test various aspects of the theory and explore additional issues.

Table of Contents

Introduction.....	1
2-AFC Testing: Preference and Perception vs. Bias and Sensitivity	4
Experiment 1: Repetition and Associative Priming.....	7
Participants.....	8
Materials.....	8
Equipment	9
Procedure	9
Results.....	13
A Note on Analyses	13
Passive Priming Results.....	14
Active Priming Results.....	15
Discussion.....	16
Responding Optimally with Unknown Sources of Evidence (ROUSE)	19
Theories of Discounting	27
ROUSE and Discounting ³	28
Setting the Value of γ	33
Setting the Vector Length.....	34
The ROUSE Model Applied to Experiment 1.....	34
Experiment 2: Orthographic/Phonemic Priming.....	36
Method.....	36
Passive Priming Results.....	39
Active Priming Results.....	40
Discussion.....	41
The ROUSE Model Applied to Experiment 2.....	42
ROUSE and Prime Similarity	44
Experiment 3: Repetition Priming and Choice Word Similarity	48
Method.....	50
Passive Priming Results.....	51
Active Priming Results.....	52
Discussion: ROUSE and Choice Word Similarity	53
The ROUSE Model Applied to Experiment 3.....	56
Experiment 4: Perceptual Benefits of Associative Priming.....	58
Method	60
Results and Discussion.....	62
Experiment 5: Associative Priming and Choice Word Similarity.....	66
Method	67
Results and Discussion.....	69

Experiment 6: Associative and Orthographic Priming.....	71
Method.....	72
Results.....	74
Discussion.....	76
ROUSE Predictions for Experiment 6	77
Experiment(s) 7: Compound Word Priming for Constituents.....	78
Experiment 7a.....	79
Method.....	79
Results and Discussion.....	82
Experiment 7b.....	84
Method.....	84
Results and Discussion.....	85
Experiment 7c.....	86
Method.....	87
Results and Discussion.....	87
Experiment 7d.....	88
Method.....	89
Results and Discussion.....	89
Experiment 7e.....	90
Method.....	91
Results and Discussion.....	93
General Discussion for Experiment 7.....	94
General Discussion.....	96
The Effects of Passive versus Active Prime Processing.....	97
Orthography versus Associations	99
"Subliminal Priming"	100
Short-term Priming Theories and Preferential Effects	102
Short-term Priming Theories and Perceptual Effects.....	104
Repetition Blindness.....	107
Long-term Repetition Priming.....	109
Is the ROUSE Theory Too Powerful?	110
Extensions of ROUSE to Other Short-term Priming Tasks	112
Conclusions.....	113
References	115
Author Note.....	121
Footnotes.....	122
Figure Captions.....	133
Curriculum Vitae.....	151

Introduction

Meyer and Schvaneveldt (1971) observed that lexical decisions were made more quickly to pairs of associated words than to pairs of unassociated words. Meyer, Schvaneveldt, and Ruddy (1974) modified the task by presenting a single "prime" word prior to lexical decision for a "target" word. In contemporary versions of this task, a prime word is presented for a duration ranging from 40 ms to several seconds, and followed by a target word to which a response must be given. Facilitation is defined as faster or more accurate responses to targets preceded by related primes than preceded by unrelated primes. Facilitation has been observed for a number of prime-target relations, including but not limited to associations (Meyer & Schvaneveldt, 1971; Evett & Humphreys, 1981; Marcel, 1983; McNamara, 1994; Perea & Gotor, 1997), mediated associations (McNamara, 1992; Mckoon & Ratcliff, 1992), semantic similarity (Perea & Gotor, 1997; McRae & Boisvert, 1998), orthographic similarity (Evett & Humphreys, 1981), phonemic similarity (Meyer, Schvaneveldt, & Ruddy, 1974), and repetitions (Evett & Humphreys, 1981; Humphreys, Besner, & Quinlan, 1988). These effects are what we refer to as short-term word priming.

In a lexical decision task participants are asked to determine as quickly and accurately as possible whether the target string of letters is a valid word; in naming, participants simply pronounce the visually presented words. In both, response time is the measure of interest. In perceptual identification target words are presented for tens of milliseconds and immediately post-masked. Participants attempt to identify the briefly flashed target

word and accuracy is the measure of interest. In these paradigms much experimentation and concern has been directed to the possibility that decision strategies (such as a tendency to respond with a word related to a prime) may affect the results.

A forced-choice technique to control decision strategies in a perceptual identification paradigm was used by Ratcliff and McKoon (1997) for long-term repetition priming (in this paradigm, the prime that is identical to the target is presented many trials prior to the test phase). In their technique two choice words (always consisting of the correct target word and an incorrect foil word) were presented soon after the brief flash of the target word. This two alternative forced choice (2-AFC) procedure proved very useful for separating perceptual and preferential aspects of long-term repetition priming. We have borrowed this technique for the sequence of short-term priming studies reported here, and utilize 2-AFC to study short-term associative, orthographic/phonemic, and repetition priming.

Within the large research effort directed toward short-term priming (e.g. see Neely, 1991 for a review of associative/semantic priming results), one major focus has been the determination of conditions leading to different amounts of facilitation (e.g. McNamara, 1992; McKoon & Ratcliff, 1992); many other studies have used short-term priming as a tool to explore various aspects of cognition. In this thesis we explore conditions producing both facilitation and decrements in performance, and ask how each should be interpreted. For example, does facilitation imply that the target words have been perceived better?

Throughout this thesis we make a distinction between priming that produces effects independent of the response to target presentation (termed ‘preferential’) and priming that produces effects by altering the perceptual processing of targets (termed ‘perceptual’).

This distinction is similar to that of Masson and Borowsky (1998) in which "contextual information" is considered separately from prime effects resulting in "perceptual encoding". With traditional word recognition tasks it is difficult to determine whether improved performance is due to preferential or perceptual mechanisms.

Whether explicitly or implicitly, preference factors could play a role. In lexical decision there could be an explicit preference to respond "word" to words related to the prime. Likewise in perceptual identification there could be an explicit preference to produce words related to the prime. Preference factors need not be explicit and might for example consist of an implicitly generated pre-activation for all prime related words. Preference effects are defined by their independence from the response of the perceptual system to the target flash; thus pre-activation is defined to be preferential if it does not alter the extraction of high or low level features in response to the target presentation. According to this definition, preference is any additive component of pre-activation. Of course interactive models of pre-activation might be invoked that would lead to better (or worse) processing of the target, perhaps through better processing of low level features, or more indirectly through increased top-down support or excitation between high level features; such interactive effects would be termed perceptual.

In our studies using perceptual identification, we gain insight into preference versus perception by manipulating the post-trial choice words; in the critical condition, both choice words are equally related to the prime. If performance in this condition is higher than that in the condition in which neither choice word is related to the prime, we argue that the additional information must have been gained from the flash of the target word. In most experiments, we also include preference conditions in which only the target or only

the foil is related to the prime. These conditions allow assessment of the direction and magnitude of preference effects.

2-AFC Testing: Preference and Perception vs. Bias and Sensitivity

Suppose that a prime is presented, followed by a brief flash of the target (e.g. SAUCE), followed by two choices (e.g. target: SAUCE vs. foil: TRAIN). The four conditions of interest are: a) **neither-primed**: both target and foil unrelated to the prime (e.g. prime = SHELF); b) **both-primed**: both target and foil related to the prime (e.g. prime = GRAVY); c) **target-primed**: the target is related to the prime but the foil is unrelated (e.g. prime = APPLE); d) **foil-primed**: the foil is related to the prime but the target is unrelated (e.g. prime = FREIGHT).

The signal detection approach (e.g. MacMillan & Creelman, 1991) assumes that at the moment of decision, there are evidence values for the choices that are selections from two evidence distributions. Forced-choice performance (e.g. $P(C)$) is inversely monotonically related to the overlap of these evidence distributions; as such performance provides a measure of sensitivity with 2-AFC testing. Bias can be thought of as the placement of a criterion for making a response to a single probe item. In the case of 2-AFC testing it is typically assumed that the alternatives are directly compared and the 'better' chosen (though a criterion could be assumed in this case as well).

The critical point is that sensitivity and bias are defined in terms of the evidence distributions accumulated over the whole task. Perceptual and preference factors, however, are defined in terms of task components; changes in evidence distributions or

criteria that are separate from evidence arising from the flash of the target itself are termed preferential, whereas changes in evidence accumulation when the target is processed are termed perceptual.

To illustrate the difference suppose that priming of both choices increases the evidence equally for target and foil, and has no other effect. We would consider this to be a preference effect. In signal detection terms both evidence distributions would shift upward, but the overlap would not change and performance (i.e. sensitivity) would not change. Additionally suppose that along with the increases in evidence, priming increases the variability of both the target and foil evidence distributions. We would still regard this to be a preference effect (i.e. no selective advantage for the target). Yet the increase in variability would increase the overlap between the distributions and reduce performance (i.e. sensitivity). Note that in cases where both choices are primed and the effect of priming equally changes target and foil distributions, it is certain that any difference from baseline is preferential; if priming enhanced or harmed target processing, there would be a selective effect upon the target evidence distribution. These ideas echo those of Norris (1995) and Masson and Borowsky (1998), who argue against equating changes in sensitivity with changes in perception.

Similar arguments can be used to describe the evidence distributions in the preference conditions, in which only the target or only the foil is primed. Importantly, priming the foil (i.e. a comparison of foil-primed to neither primed or both-primed to target-primed) can only produce a preferential effect, since until the choices appear the two conditions are identical.

It is not easy to find studies in the literature that distinguish preference from

perception. The great majority of short-term priming studies use lexical decision or naming, and obtain response time measures; these studies neither lend themselves to analyses in terms of signal detection nor unambiguously allow the separation of preference from perception. However, some perceptual identification studies with same/different responses to a single choice word used primed foil words and can be interpreted in terms of preference and perception: Johnston and Hale (1984) used a same/different procedure to study short-term repetition priming. For comparison, they used a baseline condition that contained no prime word. In a follow up study, Hochhaus and Johnston (1996) repeated the experiment with a neutral prime word baseline and obtained similar results. In both sets of experiments an analysis with repetition primed targets providing hits and repetition primed foils providing false alarms yielded reduced sensitivity as compared to the unprimed situation. In addition there was a bias in favor of repeated words. However, the sensitivity drop could be due to a decrease in perceptual encoding of targets, or to an increase in variability of preferences, or both.

Masson and Borowsky (1998, Experiments 2 and 3) used a same/different perceptual identification task to examine associative priming. In their studies, the same (targets) and different (foils) choices presented after the target flash were equally related to the prime. They found a (modest) increase in sensitivity caused by related primes, and no change in bias. Interestingly, these results held for both word primes and picture primes, even though the target presentation and the subsequent choice were words in both cases. This increase in sensitivity can be interpreted unambiguously as an improvement in target perception (certain of the studies presented in this thesis will provide a replication of these results as well as Johnston's results and place them in a larger context). Masson

and Borowsky predicted the sensitivity increase with an attractor model of priming (Masson, 1991; Masson, 1995). The theory does not assume enhanced perceptual encoding at an early stage of encoding; however, as encoding continues, prime and target presentations interact in a manner consistent with our definition of a perceptual factor.

Experiment 1: Repetition and Associative Priming

Repetition and orthographic/phonemic priming occasionally reveals deficits (Hochhaus & Johnston, 1996; O'Seaghdha & Marin, 1997; Lukatela & Turvey, 1996; Dominguez & de Vega, 1997; Humpreys, Besner, & Quinlan, 1988; Lupker & Colombo, 1994) but associative priming universally seems to produce facilitation, even when preference factors are controlled (Masson & Borowsky, 1998). We contrast repetition and associative priming in Experiment 1, in a paradigm using perceptual identification and including for each type of priming relation the four 2-AFC conditions: neither-primed, both-primed, target-primed, and foil-primed. In order to provide a both-primed condition with repetition priming, two primes are necessary and two primes were therefore presented on every trial in all conditions. In two additional conditions, both choice words were primed in mixed fashion: In one, the target was a repetition of one prime and the foil an associate of the other prime; in the other, the target was an associate of one prime and the foil a repetition of the other prime. Two versions of this study were run on separate participants, one in which participants actively processed the prime words and another in which participants passively viewed the prime words. The task for active priming required participants to determine whether the two prime words matched in term of their animacy.

Participants

There were 55 participants in the passive priming condition and 52 in the active priming condition. All participants in all experiments were native English speaking Indiana University undergraduates receiving introductory psychology course credit.

Materials

The word association norms of Nelson, McEvoy, and Schreiber (1994) were used to construct the stimulus set. These norms are based on one associative response to a prime word by each participant. As much as possible, high natural language frequency words were used (Kucera & Francis, 1967). 120 prime-associate pairs were used with average association strength of .378. Prime and associate words were 3 to 5 letters in length and could be of different lengths. All prime words could be judged as animate or inanimate (i.e. they could serve as reasonably concrete nouns). The target and foil words were drawn from the same pool of associates for all conditions, but a separate pool of words, also 3 to 5 letters in length, served as the primes for the neither-primed condition and as the unrelated primes in the target-primed and foil-primed conditions. A different pool of 4 and 5 letter words was used for the practice sessions and the threshold determination block of trials.

Randomly generated character shaped pattern masks were used to avoid pattern mask habituation. These were created by randomly selecting a position for a vertical bar within the width of a character. Then two randomly determined vertical positions were chosen on each side of the character width. These were connected to each other and to the top or

bottom of the vertical bar by separate line segments. The resulting appearance was a butterfly shape tilted to the right or the left (see Figure 1 for examples of the pattern masks). All words were displayed in capitalized, Times Roman 22-point font.

Equipment

Stimulus materials were displayed on PC monitors with presentation times synchronized to the vertical refresh. The refresh rates was 120Hz providing display increments of 8.33ms. In order to avoid phosphor decay, stimuli were displayed as black against a grey background. Subject booths were enclosed and the lighting dim to avoid eyestrain. The resulting visual contrast was close to 100%. Chin rests were used to control monitor distance. Monitor distance and font size were chosen such that target words encompassed less than 3 degrees of horizontal visual angle. All responses were collected through response boxes with 4 keys.

Procedure

Besides the neither-primed condition, the 8 priming conditions were: both repetition primed; target repetition primed; foil repetition primed; both associatively primed; target associatively primed; foil associatively primed; target repetition primed and foil associatively primed; and target associatively primed and foil repetition primed. These conditions are illustrated with examples using particular words in Table 1. Each participant received 12 trials on each primed condition and 24 trials on the neither-primed condition scattered across 120 experimental trials.

Refer to Figure 1 for the sequence of events on each experimental trial. Figure 1 is intended as a general guideline for the sequence of events within a trial as occurred in most of the experiments. The portion of Figure 1 within the dashed box only appeared for the active priming group. Each presentation sequence for perceptual identification consisted of: a fixation point for 500 ms (not shown in Figure 1); a blank screen for 500 ms (not shown in Figure 1); two prime words until response (only for active priming); two prime words for 500ms; a briefly flashed target word (of a duration determined individually for each participant); a pattern mask of duration such that the total duration of target plus mask was 500 ms; and a final display of two choice words (for 2-AFC). Setting the total target plus mask time to 500 ms equated the duration from onset of the target word (as well as offset of the prime word) until presentation of the choice words. In order to reduce word length effects, words with fewer than the maximum number of letters (5 letters being the maximum in Experiment 1) were flanked on either side by pattern mask characters. Although prime and associate could be of differing lengths, the target and foil always contained the same number of letters on a given trial as did the two prime words.

The two primes appeared above and below the center position with slightly less than 3 degrees visual angle separating them. The 2-AFC words appeared to the left and right of the center position separated by two degrees of visual angle. These choice words remained onscreen until participants responded. Participants were instructed that one of the choices would always be the flashed target word. Following their response, feedback was given before moving to the next trial.

Choice words were randomly assigned to priming condition and randomly assigned as targets or foils. Associate length and animate versus inanimate prime, were equally

assigned to the 9 different conditions. Left/right target position was counterbalanced across trials randomly. If only one of the two prime words was related to the choice words (i.e. target-primed or foil-primed), its top/down position was counterbalanced across trials randomly. For the both-primed conditions, the positions of the target related and foil related primes were similarly counterbalanced. Presentation order of the conditions was counterbalanced randomly. The conditions and number of trials per condition were such that a prime related choice word was equally likely to be target or foil. It was hoped that this would reduce any “explicit” strategy for choosing for or against prime related words. With the exception of the active priming versus passive priming manipulation, all variables were within subject. To avoid contamination from long-term repetition priming, a given word appeared only once within the experiment.

Participants received 16 trials of practice on perceptual identification. In the active priming condition, perceptual identification practice was preceded by 16 trials of practice on the animacy matching task in isolation; all subsequent trials included both animacy matching and perceptual identification. For the animacy matching task, participants were instructed to press a key labeled "match" if the two prime words matched in animacy and otherwise press a key labeled "mismatch". During practice, they immediately received feedback on their animacy judgments. Thereafter, every eighth trial they received cumulative feedback for the last 8 trials. Following their animacy decision, prime word(s) remained on the screen for an additional 500ms before being replaced by the flash of the target word (see Figure 1). This was done so participants would not miss the target flash while responding to the prime word(s). As compared to the display for animacy matching, the prime words switched locations and were displayed in bold face, during this 500ms. In

the passive priming condition, prime words appeared in bold face for 500ms (hence the 500ms prior to target presentation where identical in both conditions). For passive priming, participants were instructed that prime words were a warning to prepare for the flash of the target word.

Following initial practice trials, participants were presented with a block of 72 perceptual identification trials. The purpose of this block was to find the duration of target flash at which performance was 75%. Appropriate durations averaged about 50 ms, although there were large individual differences, with times ranging from 25ms to 117ms. A staircase method was used to find the appropriate target duration during this threshold determination block. Participants were fully informed about the procedure. The words for the threshold determination block and practice trials were randomly selected (i.e. neither-primed). Prime relatedness was not introduced until the experimental trials.

 Insert Figure 1 about here

 Insert Table 1 about here

 Insert Figure 2 about here

Results

A Note on Analyses

One source of evidence concerning possible effects of priming upon target perception is obtained from subtraction of neither-primed $P(C)$ from both-primed $P(C)$, assessed statistically by an appropriate t-test. Improved performance in the both-primed condition provides evidence of a perceptual enhancement. The idea is that preference effects are equated if both choices are primed, and that priming might increase variability of evidence but would be unlikely to decrease variability. An increase in performance therefore provides relatively unambiguous evidence for a beneficial effect of priming upon target perception. If a deficit is observed it is evidence for the existence of increased variability with priming (preferential variability) and/or perceptual inhibition; such a result leaves open the possibility that perceptual enhancement exists, as long as the (negative) effect of preferential variability is large enough to overcome the perceptual enhancement.

A comparison of neither-primed (a) to foil-primed (d) involves equal effects upon target perception, since in both cases the prime is unrelated to the target (i.e. in both cases the displays are identical up until the 2-AFC). Therefore the difference between performance in these conditions is as indicator of preference effects. Similarly the comparison of target primed (c) to both primed (b) involves equal effects upon target perception since in both cases the prime is related to the target. Therefore the difference between these is again an indicator of preference effects. One of our analyses combines these as $(a-d) + (c-b)$ and tests the result with an appropriate F test. It should be noted that a positive result of this test indicates the presence of preference effects, but the failure of the test does not disprove preference since preference and preferential variability can

potentially counteract one another. For example, if there's a slight preference against a primed foil, the corresponding increase in performance can be offset by the decrease in performance associated with increased variability.

A final analysis assesses whether the preference is in favor of or against prime related words (i.e. the direction of preference): The foil-primed $P(C)$ is subtracted from the target-primed $P(C)$, and an appropriate t-test used to carry out statistical analysis. Although this difference includes any perceptual effects of target priming, it will be useful in combination with the first two tests, and particularly in comparisons of the active and passive priming conditions.

Passive Priming Results

The upper panel of Figure 2 and the passive probability correct column of Table 1 show the accuracy results for the various conditions. (The predictions shown in Figure 2 will be discussed in the next section of this thesis.) There was a difference between repetition priming and associative priming, $F(1,54) = 6.89$, $p < .025$ that interacted with the four priming conditions, $F(3,52) = 7.39$, $p < .001$. This difference reflects the fact that there was a deficit in the both-primed condition for repetition priming only. In addition, participants tended to choose the related choice word, whether the relation was associative or repetition, although this effect was larger for repetitions.

For repetition priming, preferential variability or a perceptual deficit or both was a factor: Performance in the both-primed condition was worse than the target-primed condition, $t(54) = 2.92$, $p < .005$. In addition there was a large overall preference effect,

$F(1,54) = 43.18$, $p < .001$; this consisted of a preference to choose the word that repeated a prime, $t(54) = 6.0$, $p < .001$.

For associative priming, performance in the both-primed condition was not different than the target-primed condition, $t(54) = .97$, $p = .34$. There was an overall preference effect, $F(1,54) = 7.73$, $p < .01$, to choose the word that was an associate of the prime, $t(54) = 6.0$, $p < .001$.

When the target was repetition-primed and the foil associatively-primed, performance was higher than when the foil was repetition-primed and the target was associatively-primed, $t(54) = 3.16$, $p < .0025$. In other words participants tended to choose the repeated word even if the alternative word was associatively primed.

Active Priming Results

Participants took an average of 3173 ms with a $P(C)$ of .732 in the animacy matching task, suggesting this was a difficult task involving considerable processing of the primes.

The lower panel of Figure 2 and the active probability correct column of Table 1 show the accuracy results for the various conditions. There were differences between associative and repetition priming, $F(1,51) = 19.95$, $p < .001$, that interacted with the four basic priming conditions, $F(3,49) = 10.99$, $p < .001$. The difference was due to a deficit in both-primed performance (i.e. preferential variability or perceptual deficit), and a difference between the target-primed and foil-primed conditions, each of which only occurred for repetition priming. Surprisingly, participants tended to choose the choice word that was **not** a repetition of a prime. Thus the direction of preference was opposite to that seen with passive priming

Within repetition priming conditions, performance in the both-primed condition was worse than the neither-primed condition, $t(51) = 4.62$, $p < .001$, suggesting preferential variability outweighed any improvement in target perception caused by prime repetition. Although there was no preference effect according to the combined measure, $F(1,51) = 2.41$, $p = .13$, the individual comparison of the both-primed condition to the target-primed condition, $t(51) = 2.97$, $p < .0025$, revealed evidence of a preference effect. In addition the comparison of target-primed to foil-primed shows that the preference was against repeated words (target-primed lower than foil-primed, $t(51) = 2.14$, $p < .025$). Within the associative priming conditions, there were no differences across the four basic conditions, $F(3,51) = 1.04$, $p = .39$.

When the target was repetition-primed and the foil associatively-primed, performance was lower than when the foil was repetition-primed and target associatively-primed, although this difference did not reach significance, $t(51) = .98$, $p = .17$. In other words there may have been a slight preference to choose the non-repeated word even if that word had been associatively primed.

Discussion

The first noteworthy result is found in the repetition condition, for both active and passive prime processing: There was a substantial deficit in the both-primed condition compared with the neither-primed condition. There are two obvious hypotheses, either or both of which could be true: 1) the primes produce a deficit in perceptual processing of the target; 2) the increase in variability of evidence induced by the prime (i.e. preferential variability) outweighs any improvement in perceptual processing of the target (if there is

any such gain). The ROUSE model presented in the next section largely explains the results in terms of explanation 2.

The second noteworthy result is the switch from a preference for repeated words with passive priming to a preference against repeated words with active priming. Using traditional priming tasks, which presumably include both preferential and perceptual effects, decreased performance with orthographic/phonemic priming is occasionally observed (O'Seaghdha & Marin, 1997; Lukatela & Turvey, 1996; Dominguez & de Vega, 1997; Lupker & Colombo, 1994). Presumably, whatever is responsible for these deficits should apply at least as strongly to the case of repetition priming. Typical explanations for these deficits appeal to lexical suppression (e.g. Lupker & Colombo, 1994) or phonological competition (e.g. O'Seaghdha & Marin, 1997). The ROUSE theory provides a unique perspective on these occasionally observed deficits by proposing they are the product of a preference against prime related words. The theory predicts negative preference to be somewhat unreliable and only elicited in circumstances prompting participants to more fully (i.e. actively) process the primes. In situations where primes are processed to a lesser degree (i.e. passively), the more commonly observed facilitation with priming will result.

Repetition priming is rarely studied, except with sub threshold prime presentations, due to concerns of strategic responding. However, within a Rapid Serial Visual Presentation (RSVP) sequence of words, the effect of presenting a word upon its later re-presentation has been studied. This paradigm led to the observation known as 'repetition blindness' (Kanwisher, 1996). In a typical repetition blindness experiment, participants fail to notice the second presentation of a word. Sometimes cited as an example of repetition

blindness, Johnston and Hale (1984) and Hochhaus and Johnston (1996) also found repetition deficits. These experiments are unique in that, similar to our Experiment 1, preference factors were controlled. In these studies participants were not required to respond to the prime word, analogous to our passive priming condition. These studies also found a bias in favor of repeated words (similar to our passive priming results). Any simple interpretation of such bias, however, must explain the reversal of this tendency in our active priming condition.

Unlike the results of Masson and Borowsky (1998) we do not find an associative enhancement when preference factors are controlled. However, associative effects in our studies might be weakened by the use of two primes, a hypothesis shown to be correct in later experiments. More generally, our repetition and associative priming results provide strong evidence that priming produces preference effects, but no direct evidence that priming produces perceptual changes in target processing. Thus the ROUSE model presented in the next section is a model of preference effects.

With passive priming, the tendency to choose prime related words occurs in associative priming as well as repetition priming (i.e. a positive preference). With active priming, this tendency is reversed for repetition priming. Yet for both active and passive priming there is a deficit in the both repetition primed condition. This pattern of results presents a complex set of interactions that seems at first glance to defy simple explanation. We shall see next that this assessment is incorrect, because a rather simple Bayesian model of decision making for this task provides a coherent account of the pattern of results. The theory accounts for the data assuming only preferential factors are involved in priming.

Responding Optimally with Unknown Sources of Evidence (ROUSE)

Our theory supposes preference effects occur because some features of the prime are carried over and confused with features extracted at the time of the flashed target. A choice word containing such features thereby tends to be favored. One can perhaps think of this activation of choice word features as a kind of pre-activation. However, unlike the pre-activation in spreading activation theories (Collins & Loftus, 1975; Anderson, 1983), we assume pre-activation only applies to shared features between prime and choice word and as such only part of the representation might be affected (to the degree that features are shared). This pre-activation does not affect the process of target activation and is therefore preferential in nature. To be more precise, the theory holds that the primes, the flash of the target, and general visual noise are all independent sources of feature activation. A preference for prime related words arises due to a failure to distinguish the various sources of activation (see Johnson, Hashtroudi, & Lindsay, 1993 for a review of source monitoring phenomena).

We shall see that the variability in activation of choice-word features by the primes is the basis for the deficits observed in the both-primed repetition conditions. In addition, we shall see that slight inaccuracies in the attempt to take into account the uncertainty in the origin of activated features leads to the different directions of preference effects:

Responding with prime related words in passive conditions and the opposite in active

conditions. We call the proposed theory Responding Optimally with Unknown Sources of Evidence, or ROUSE for short.

In the present task the participant must choose between two words. Let us assume that a vector of feature values represents each word. We used a vector of length 20 in our simulations (though Figure 3 uses 10 features to reduce clutter). Assume further that each feature is binary, either being ON or OFF. At the start of a trial all features are OFF. The features are turned ON by three sources of evidence (see Figure 3)¹:

- 1) **The target flash.** With probability β , each feature in either choice that is in the target is activated (β depends on flash duration).
- 2) **Visual noise.** With probability γ any feature in either choice word is turned on.
- 3) **The primes.** With probability α any feature that is shared between a prime and either choice word is turned on².

Similarity: In some conditions two words are allowed or assumed to share some but not all features. When this is the case, the probability that a feature in one of the words will be identical to the corresponding feature in the other word is a parameter, ρ . For example, in Experiment 1 we assume that the associative priming conditions involves partial feature sharing between primes and primed choices, and therefore let ρ be the probability that a feature of an associatively primed choice word is shared with the prime. For simplicity we assume that unrelated words share no features ($\rho = 0$), and repetitions share all features ($\rho = 1$). Associative and similarity relations involve some proportion of shared features that must be estimated in order to apply the model. It should be emphasized that the participant need not pay heed to the value of ρ . Since the prime words and choice words are

all known on a given trial, we assume that it is clear to the participant which features overlap and which do not.

Note: We assume a feature common to the two choice words is ignored in the decision process. This assumption follows from the math if shared features exist in the same state for both choices (i.e. both ON or both OFF).

 Insert Figure 3 about here

Each of the three activation probabilities (α , β , and γ) represent the joint probability that the source has turned on a feature and that the feature has remained on until the time of the decision process. Because the features produced by prime activation will eventually decay with time or intervening events, they will cease to cause confusions. ROUSE is a model of short-term priming in that priming arises through activation and is therefore not expected to exist over extended durations.

Source confusion naturally leads to a preference for prime related words. In the decision process we suppose participants attempt to remove the irrelevant prime activation and that the extent of this removal leads to the various preference results. At the time of decision the features that are turned ON and OFF in each choice word are assessed and an optimal decision realized by calculating the likelihood ratio for each choice word. The word with the larger likelihood ratio is chosen. The calculations require appropriate estimates of the three activation parameters (the parameter ρ does not require such an estimate). We assume to start with that the estimates of β and γ are accurate. (If we let the estimates of β and γ differ somewhat from their true values, overall performance changes,

but the qualitative pattern of predictions across the priming conditions does not change significantly). Most important, we assume that the estimate of prime feature activation (labeled α') is sometimes too low ($\alpha' < \alpha$ in the passive prime condition), and sometimes too high ($\alpha' > \alpha$ in the active prime condition).

This close-to-optimal Bayesian decision process discounts the force of evidence from features that are ON but also known to exist in the primes; such discounting is appropriate since the primes rather than the target might have been the source of activation. By way of example, suppose that the features are letters, that a T is perceived in the first letter position, that the choice words are TOWN and SEAM, and that neither prime word has a T in first position. The perceived T provides good evidence that the flashed word was TOWN since the T could only have come from the flash or visual noise. On the other hand suppose that one prime began with a T; in this case the perceived T could have come from the prime so the evidence in favor of the choice TOWN should be somewhat mistrusted. In principle, such discounting would apply only to choice features that the participant knows are also present in the primes; in the present task the primes are presented well above threshold, even in the passive priming condition, so we assume that it is known which choice word features are also in the primes. (Below, we discuss similar notions of discounting that have been proposed for priming within the social cognition field).

The appropriate level of evidence follows from a Bayesian calculation. The odds for the target over the foil is given in Equation 1 (**A** refers to the target word and **B** refers to the foil word). Assuming equal priors, as appropriate for our experimental design, the normative decision is to choose the target **A** if the odds is greater than one, and to choose

the foil **B** if the odds is less than one. If the odds is exactly equal to one, the target is chosen with probability .5.

$$\Phi\left(\frac{\mathbf{A}}{\mathbf{B}}\right) = \frac{p\langle \mathbf{A}, \mathbf{B} \text{ activation pattern} | \mathbf{A} \text{ is target, } \mathbf{B} \text{ is foil} \rangle}{p\langle \mathbf{A}, \mathbf{B} \text{ activation pattern} | \mathbf{B} \text{ is target, } \mathbf{A} \text{ is foil} \rangle} \quad (1)$$

In an optimal treatment of feature activation, an active feature provides evidence in favor of a choice word as target, but an inactive feature provides evidence against a choice word as target. Since both active and inactive features provide useful evidence, all features in both choice words contribute to the likelihood calculation, excepting only features that are common to both choices. Assuming each feature yields an independent source of evidence and breaking Equation 1 into two separate products for the features of each choice word leads to:

$$\Phi\left(\frac{\mathbf{A}}{\mathbf{B}}\right) = \frac{\prod_{i=1}^N \frac{p\langle V(\mathbf{A}_i) | \mathbf{A} \text{ is target} \rangle}{p\langle V(\mathbf{A}_i) | \mathbf{A} \text{ is foil} \rangle}}{\prod_{j=1}^N \frac{p\langle V(\mathbf{B}_j) | \mathbf{B} \text{ is target} \rangle}{p\langle V(\mathbf{B}_j) | \mathbf{B} \text{ is foil} \rangle}} \quad (2)$$

In Equation 2, $V(\mathbf{A}_i)$ represents the value of the i -th feature of **A**, and takes on one of two values, which denote the ON and OFF states of activation; similarly for $V(\mathbf{B}_i)$. The product in the numerator is termed the likelihood ratio for the target word, and is based on evidence coming from the N features of the target **A**; the product in the denominator is termed the likelihood ratio for the foil word, and is based on evidence coming from the N features of the foil **B**. Thus the choice of **A** if the odds in Equation 1 is greater than one is equivalent to the choice of **A** if its likelihood ratio in Equation 2 is greater than that for **B**.

There are only four possible evidence ratios that could appear in the product terms found in Equation 2. These are shown in Figure 4, depending upon whether a feature is

ON or OFF and whether a feature is known to have appeared in the prime(s) (i.e. whether a prime is a potential source of activation). For instance consider an OFF feature that did not appear in the prime(s). Assuming that the feature is part of the target, the target flash and/or noise could have been a source of activation. Since the feature is OFF, both of these sources must have failed so the probability is $(1-\gamma)(1-\beta)$. Assuming that the feature is not part of the target, only noise is a potential source of activation, so the feature is OFF with probability $(1-\gamma)$. This leads to the ratio seen in the upper left cell of Figure 4 (which is equal to $1-\beta$). A related calculation produces the same result, $1-\beta$, in the lower left cell of Figure 4. In other words, prime and noise activation are irrelevant to the evidence provided by OFF features; only target activation matters.

 Insert Figure 4 about here

The lower right cell in Figure 4 gives the evidence for an ON feature that appeared in the prime(s); such a feature is termed ‘discounted’ because its evidence ratio is less (i.e. closer to one) than if it had not appeared in the prime(s) (the term in the upper right cell). It is the estimate of prime activation, α' , that will determine the level of discounting.

The relative size of prime activation, α , compared to the estimate of prime activation, α' , produces the direction of the preference: e.g. whether the target-primed condition is better than the foil-primed condition. If participants are optimal (i.e. $\alpha' = \alpha$) and the number of diagnostic features turned on by each of the sources of activation is sufficiently large, feature evidence from primes will be discounted properly, and there will not be a positive or negative tendency to choose a word related to the prime. (Note however, that

the overall performance of these preference conditions taken together is still predicted to be lower than performance in the neither-primed conditions, because variability will exist in the number of prime-activated features; see the explanation in the second paragraph below).

If participants are conservative and overestimate the effect of the prime (i.e. $\alpha' > \alpha$), and if the number of diagnostic features turned on by each of the sources of activation is sufficiently large, words that are *not* related to the prime(s) will tend to be chosen. As we shall see, such a situation is what we assume exists with active priming. If participants are less aware of the prime(s) as a potential source of activation, they may underestimate the effect of the prime (i.e. $\alpha' < \alpha$). If so, they will tend to choose prime related words. This is the situation we assume holds with passive priming. A pictorial explication of these arguments concerning preference effects is given in the section below headed '*ROUSE and Discounting*'.

Next consider predicted performance in the both-primed conditions. On average the number of ON features that are shared with the primes will be the same for the two choice words (i.e. preference is controlled). However, there will be variability in these numbers, so that sometimes one and sometimes the other choice word will be favored, purely by chance. This chance process adds noise to the decision, decreasing predicted performance compared with the neither-primed condition. The size of this variability effect will depend on the values of α and ρ , with increasing variability for larger values of either parameter. This means the deficit will be largest for repetition priming, for which all features are shared.

This effect of variability of evidence arising from prime activation is best described as preferential because it occurs even when there is no change in the perception of the target. However, primes have a second effect upon predicted performance, an effect that might be described either as perceptual or preferential, depending on one's perspective: Features that are turned on by the prime are unavailable to be turned on by the target. This is an example of performance being harmed by a prime, due to 'blocking'. Considering the situation from the point of view of the features that are eventually ON, one can describe the harm as perceptual, because perception is blocked from producing distinguishing evidence. For example, if the parameter α equals 1.0, then in the both-primed condition every feature of both choice words is turned on by the primes; now the flash of the target provides no information, and performance is at chance. One could argue that no perception has occurred. On the other hand, one could argue that perception is unaffected by the primes, but that the evidence provided by perception is overwritten by the prime features; in this case, one might prefer to describe the blocking effect as preferential. Fortunately, this debate in the context of the present studies is of little consequence: It turns out that the fit of the model to the data resulted in quite low estimated values of α and β (and γ). With low values for these parameters, a feature turned on by a prime is rarely also turned on by the target flash, so the effect of blocking turns out in practice to be of negligible importance; this being so, it is not critical to decide whether blocking should be thought of as a perceptual or preferential effect. In particular the both-primed deficit predicted with ROUSE is almost wholly due to preferential variability rather than blocking.

The associative case is similar to the repetition case, differing only in having relatively few features shared between prime and target, or between prime and foil ($p < 1$). This change lessens the effects of priming generally. It should be noted again that although we need to estimate the value of p to fit the model, the participant need not estimate the value of p : Both the prime features and choice word features are available to the participant on each trial (assuming the participant pays attention to the above threshold primes), so the shared features are identifiable and do not have to be estimated by the participant.

Theories of Discounting

The idea of discounting evidence is far from unique to the present treatment. A particularly relevant example arises in the area of evaluative priming. Similar to the preference reversal observed in Experiment 1, social psychologists have observed priming reversals in evaluative judgments. Lombardi, Higgins, and Bargh (1987) had participants construct sentences using synonyms of ‘persistent’ versus ‘stubborn’. Following this, participants gave a one-word label to a neutral description of a target person. Participants who later remembered their constructed sentences tended to use a label the opposite of the synonyms they received (i.e. if their sentence included ‘determined’ they rated the target person as ‘stubborn’). Conversely, participants who could not remember their sentences tended to use a label similar to the synonyms.

Experiments such as this observing ‘contrast’ (i.e. a preference against prime related words) and ‘assimilation’ (i.e. a preference for prime related words) priming have been explained by theories supposing participants might or might not attempt to remove the influence of prime items (Martin, 1986; Schwarz and Bless, 1992), or more generally to

remove what is believed to be a contaminating influence of certain noticed mental events (e.g. Wegener & Petty, 1995). In assimilation priming the effect of the prime is realized in full whereas in contrast priming the effect of the prime is removed resulting in preference against prime related words. One theory of discounting holds that source similarity between prime and target is of particular importance (Musseweiler & Neumann, in press). With externally presented targets, it is predicted that externally presented primes lend themselves to discounting whereas internally generated primes are less likely to be discounted. Indeed, Musseweiler and Neumann found that internally generated primes (e.g. antonym generation task) lead to assimilative priming whereas a simple presentation of these same primes lead to contrast priming.

We note, however, that these theories of discounting tend to produce quite explicit and broadband effects. We shall see in the remaining studies that the discounting mechanism at work in the present situation is much more subtle than entailed by such theories.

ROUSE and Discounting³

To explicate the discounting mechanism in ROUSE that removes or even reverses preference, we present a Venn diagram in Figure 5 for the conditions in which either the target-only or foil-only is primed. In this figure we illustrate how the direction of preference (or lack thereof) arises from two offsetting factors. We do not use this figure to explain the both-primed deficits in performance caused by priming, because these are due to variability, and the figure depicts averages. Figure 5 illustrates the evidence situation for the target and foil when similarity between prime and one of the choices is intermediate

(e.g. associative or orthographic/phonemic priming, $0 < \rho < 1$). The left-hand set of panels shows the situation without discounting (i.e. $\alpha' = 0$) and the right-hand set of panels shows the situation with optimal discounting (i.e. $\alpha' = \alpha$).

In general, features can provide one of three levels of evidence. Features that are OFF provide evidence against a choice word regardless of prime matching (as seen in Figure 4); these are represented as the black regions in Figure 5. Features that are ON provide a full force of evidence in favor of a choice word if they do not appear in the prime(s); the white regions in Figure 5 represent this situation. Lastly, ON features that appeared in the prime(s) provide a reduced level of evidence in favor of a choice word (since the evidence is discounted); this situation is represented by the grey regions in Figure 5: The higher is the level of discounting (i.e. the higher is α'), the darker is the shade of grey, and the lower is the evidence provided by such features. In the left-hand set of panels there is no discounting so all active features are displayed in white. Performance (i.e. $P(c)$) corresponds to the product of the evidence from each feature for the target as compared to same product for the foil (see Equation 2); in Figure 5 this may be thought of as the average 'lightness' of the target evidence as compared to the average 'lightness' of the corresponding foil evidence.

Without discounting (i.e. $\alpha' = 0$) it is clearly seen that performance is bolstered in the target-primed condition: There is more white for the target evidence panel of the target-primed condition due to the presence of prime features, as indicated by the region labeled 'PRIME(S)'; for the foil-primed condition this region switches to the foil evidence panel.

With discounting (i.e. $\alpha' > 0$), the relation of target-primed $P(c)$ to foil-primed $P(c)$ depends on the balance of two offsetting forces: 1) *Prime Activation*: The prime tends to

turn otherwise OFF features ON, leading to a preference to choose the related choice word; 2) *Discounting Target/Noise Activation*: The prime causes features that are turned ON by the target flash, and by noise, to be discounted, leading to a preference to choose the unrelated choice word. *Prime Activation* corresponds to the grey regions labeled 'PRIME(S)'. Even though these ON features are discounted they still provide a preference for prime related words since discounted evidence is still better evidence than OFF features. *Discounting Target/Noise Activation* corresponds to the small grey circles within the 'TARGET' and 'NOISE' regions. These are features activated by the target and/or noise that are discounted since they also exist in the primes. Without discounting these features would provide strong ON evidence but with discounting they are mistrusted since the primes might have been the source of activation. This leads to a loss of evidence in favor of the primed alternative as compared to situation without discounting. Therefore these features provide a preference against primed words.

Insert Figure 5 about here

If the right degree of 'greyness' is chosen (usually corresponding to $\alpha' = \alpha$, which we have termed 'optimal') the tradeoff between extra *Prime Activation* evidence and the loss of evidence due to *Discounting Target/Noise Activation* results in equal performance for the target-primed and foil-primed conditions. If similarity between prime and choice word is increased, then both the size of the grey circle labeled 'PRIME(S)' and the number and

density of the small grey circles increase, typically maintaining the balance between the two opposing forces. For example, for repetition priming ($\rho = 1$) all the activated features (circles) in the upper left and lower right panels in the right half of Figure 5 turn grey; simultaneously, the size of the circles labeled 'PRIMES(S)' increases, maintaining the overall degree of 'lightness' of these panels.

If discounting is insufficient (i.e. $\alpha' < \alpha$), such as is the case with passive priming, the grey is too light, and the situation approaches that in the left hand part of Figure 5: target-primed performance better than foil-primed. If discounting is excessive (i.e. $\alpha' > \alpha$), such as is the case with active priming, the grey is too dark. This state of affairs usually results in target-primed worse than foil-primed, but there are important exceptions. If the number of features is sufficiently large and the activation parameters are not too small, discounting can always lead to a preference against prime related words provided α' is enough larger than α . However, since the model includes a finite number of features, non-linearities are introduced: with a small number of features or small activation parameters, the effect of *Discounting Target/Noise Activation* can be diminished or even lost resulting in a preference for prime related words regardless of the level of discounting. These breakdowns in preference removal are encountered in subsequent experiments and the nature of each breakdown is considered in the discussions: '*ROUSE and Prime Similarity*' and '*ROUSE and Choice Word Similarity*'.

Distributions of Odds

Insert Figure 6 and Table 2 about here

The Venn diagram represents an attempt to provide insight into the tradeoff of factors that govern performance by depicting mean numbers of different types of activated features, and their associated evidence values. An alternative and more precise way to illustrate the working of the model is through graphs of the distributions of odds. These distributions depend on the probabilities of obtaining various numbers of ON and OFF features in the different conditions. The distribution of the odds is highly skewed (because of the multiplication of probabilities), so for clarity we plot the distribution of the log-odds. In viewing the graphs in Figure 6, recall that a correct decision is made if the odds of Equation 1 is greater than 1.0, equivalent to the log-odds being greater than zero. The distributions are also quite discrete because only certain combinations of ON and OFF features are at all likely. Figure 6 shows distributions for the four basic conditions (neither primed, both primed, target primed, foil primed), for repetition priming ($\rho = 1$) with typical activation parameters ($\beta = .05$, $\gamma = .02$, and $\alpha = .1$), for three values of α' : 0, .05, and .3 (corresponding to no discounting, discounting that is too small in light of the actual value of α , and discounting that is too large in light of the actual value of α). Table 2 provides some summary statistics concerning these distributions.

Looking first at the case $\alpha' = 0$, we see that priming both alternatives adds noise and causes the distribution to spread in comparison with the neither-primed condition; although the mode remains at the same position, the extra variability causes more of the distribution to fall below zero, reducing performance. Priming only the target or only the foil adds some variability, but the primary effect is to shift the log-odds in favor of the primed choice. The case $\alpha' = .05$ includes discounting, which lessens the evidence value for the features that are in the primes, squeezing the log-odds toward zero in all three

conditions with priming (the neither-primed condition is only displayed once since it is the same regardless of α'). Performance overall is as expected, with a drop for both-primed and a preference for the primed alternative (as compared to $\alpha' = 0$, the preference is diminished). The case $\alpha' = .3$ discounts much more strongly, shrinking the distributions severely, but still producing a both-primed deficit. The fact that too much discounting occurs in this case reverses the direction of preference; the unprimed choice tends to be chosen, improving performance when the foil is primed, and harming performance when the target is primed.

For the target-primed and foil-primed conditions with discounting, the ON features of the unprimed alternative provide strong evidence, whereas the ON features of the primed alternative are discounted and this results in extra variation producing more discrete points in the distributions. For the both-primed condition all ON features are discounted in both alternatives resulting in the same degree of variation (i.e. number of spikes in the distribution) as seen in the both-primed condition with $\alpha' = 0$; correspondingly the both-primed deficit is exactly the same for all values of α' . This fact points out the significance of the both-primed condition: When both alternatives are primed, performance is unaffected by the direction of preference (i.e. $P(c)$ independent of α'). Nevertheless performance is affected by preferential variability and the same degree of variability (not variance) applies regardless of the extent of discounting.

Setting the Value of g

We discovered when fitting ROUSE to the data from the various studies in this thesis that the value of γ , the contribution of visual noise, was usually estimated to be quite low

(and needed to be low to produce an acceptable fit), and that the fits were seldom harmed if the value was set to some fixed value close to zero. We therefore simply set the value of γ to .02 throughout this thesis, rather than estimate it as a separate parameter.

Setting the Vector Length

Generally speaking we found that the vector length (above a certain minimum below which preference removal breaks down – see ‘*ROUSE and Choice Word Similarity*’ for a discussion of this breakdown) was a scaling factor whose value would be more important for determining the length of time needed to carry out simulations than determining the pattern of predictions: Longer lengths (such as 100, say) would produce similar predictions to those for shorter lengths, once suitable modifications were made to the values of the other parameters. A length of 20 was long enough to enable preference removal but was short enough to allow parameter fitting to be carried out in reasonable time.

The ROUSE Model Applied to Experiment 1

 Insert Table 3 about here

To fit the ROUSE model to the results of either the active or passive group in Experiment 1 requires estimation of the four parameters: α , α' , β , and ρ (γ was set to .02, not estimated). These four parameters are separately estimated for the two groups of participants. Averaging across 20,000 simulations to obtain predictions for a given set of

parameters, the parameters were assigned values that minimized the error between predictions and observations. The error measure used was the sum of the chi-squares from each condition where chi-square was calculated using the normal theory maximum likelihood method (see Curran, West, & Finch, 1996, for a comparison of various methods for calculating χ^2). Table 3 gives the parameter values that produced a best fit to the data from Experiment 1. The resultant predictions are given in Figure 2 as the dots on each bar.

It is clear that this rather simple model manages to capture the essentials of the data, including the complex pattern of interactions that we commented upon earlier. There are a few things that can be said about the parameter estimates. The estimate of α (i.e. α') is lower than α for passive priming, and higher than α for active priming, as needed to reverse the direction of preference (the sharp eyed reader will note that the associative predictions for the active group show a slight preference for targets, the opposite direction of preference predicted for repetition for the active group; this issue will be discussed following Experiment 2). The most surprising estimates are those for ρ , since the value is so much smaller for the active group. This passive/active difference came about because the active priming group exhibited no associative priming effects, requiring a low estimate of ρ . It is possible that the active priming instructions prompted participants to think about the animacy characteristics of the two prime words, disrupting the natural tendency to attend to other sorts of features that are used to generate associates in production tasks. We explore this issue further in a later experiment that uses active priming but only a single prime.

Experiment 2: Orthographic/Phonemic Priming

The results of Experiment 1 and the success of the ROUSE model suggest that some features from the prime words are confused with features from the flashed target. The difference between repetition and associative priming in Experiment 1 is explained by ROUSE in terms of the number of features in the primes that are shared with the features in the choice words (i.e. the number of confusable features). Presumably, primes and associated choice words only share semantic features whereas repetitions share semantic, orthographic, and phonemic features. In Experiment 1 preference effects were greatly reduced for associative priming implying orthographic/phonemic features are primarily responsible for the preference effects with repetition priming. Therefore in Experiment 2 we test this idea through conditions in which orthographic/phonemic similarity between primes and choice words is retained, but semantic similarity is removed. If, as predicted by ROUSE, shared features generally determine preference effects and, as implied by Experiment 1, the shared features in repetition priming are primarily orthographic and phonemic, then highly similar orthographic/phonemic primes should lead to a pattern of results similar to those obtained in Experiment 1 with repetition priming.

Method

There were 52 participants in the passive priming condition and 56 in the active priming condition.

Four categories of prime-choice word pairs were created each consisting of 48 five-letter words and 48 four-letter words. The pairs in three of the categories were orthographically very similar since 4 out of 5 or 3 out of 4 letters were identical and in the same position within the letter string. The remaining category was used for *repetition* priming. The orthographically similar categories were further subdivided by degree of phonemic and semantic similarity. The first category, labeled *orthographic*, was not semantically similar and less phonemically similar and included pairs such as ANGEL→ANGER. The second category, labeled *orthographic and phonemic*, was more phonemically similar but not semantically similar and included pairs such as, ALTAR→ALTER. The third category, labeled *orthographic and semantic*, was semantically similar (although not necessarily phonemically similar) and included pairs such as AWAKE→AWOKE. The four basic priming conditions were run for each of the four categories of words resulting in a total of 16 conditions. These conditions are illustrated with examples using particular words in Table 4. Each participant received 12 trials on each of these conditions.

Insert Table 4 about here

Assuming orthographic priming produces large preference effects, it is unclear whether visual or abstract orthographic features are responsible. Therefore, letter case between the primes and choice words was manipulated. For the passive priming condition, the *orthographic* (e.g. ANGEL→ANGER) and *orthographic and phonemic* (e.g. ALTAR→ALTER) pairs were combined into one category. This larger category was then

split randomly in half for each subject. For one half the pairs the primes were presented in lower case and for the other half, the primes were presented in upper case. Target flash and choice words were always presented in lower case, hence there was a case switch between primes and target for the latter group of words. As seen in Table 4, switching case did not lessen preference effects and therefore the *orthographic* and *orthographic and phonemic* categories were kept separate and all words presented in lower case in the active priming condition.

In order to ascertain the generality of the passive/active difference, a different active priming task was used. This task was to determine if the two prime words could serve the same part of speech. If the two prime words could serve the same part of speech, participants were instructed to press a key labeled "match" and otherwise press a key labeled "mismatch". No feedback was given on this task and accuracy was not calculated.

Experiment 1 used a separate pool of words for unrelated prime words. In order to control against any confounds introduced by using different primes, Experiment 2 and all subsequent experiments created unrelated prime words through a re-pairing technique. Conditions priming only one or neither choice word were created by randomly re-pairing prime and choice words such that across participants the same prime words were used in each of the priming conditions. In other words, primes and their related choice words could be presented in an intact or rearranged form. For instance, one participant might receive the intact pairs ANGEL priming ANGER as well as CHOIR priming CHAIR while another participant might receive the rearranged pairs ANGEL priming CHAIR and CHOIR priming ANGER⁴.

All other procedures were the same as Experiment 1.

Passive Priming Results

The passive priming results are shown in Table 4. There were differences across the four types of primes, $F(3,49) = 8.44$, $p < .001$, and these interacted with the four priming conditions, $F(9,43) = 5.39$, $p < .001$. The repetition priming results from Experiment 1 are replicated for passive priming: a large both-primed deficit, $t(51) = 4.94$, $p < .001$, an overall preference effect, $F(1,51) = 38.12$, $p < .001$, and a preference to choose repeated words, $t(51) = 2.82$, $p < .005$.

Three types of orthographic priming were used: two types with or without a case change and a third type using the *orthographic and semantic* word pairs. There were no differences among the three types of orthographic priming, $F(2,50) = .94$, $p = .40$, and no interaction between these types and the four priming conditions, $F(6,46) = 2.17$, $p = .06$. Separately analyzing the two types of orthographic priming between which the case of the primes varied, there was no effect of switching the case of the prime, $F(1,51) = .75$, $p = .39$, and no interaction between switching the case of the prime and the four priming conditions, $F(3,49) = 1.59$, $p = .21$ (see Table 4). This result strongly suggests abstract orthography rather than visual similarity is crucial to prime interference, and more generally that abstract orthography is the level at which most features are compared and utilized in the present tasks. Because the orthographic priming conditions did not differ, they are collapsed and the results depicted in Figure 7.

Insert Figure 7 about here

Analyzing the collapsed orthographic conditions, performance in the both-primed condition was no different than the neither-primed condition, $t(54) = .90$, $p = .37$. There was however an overall preferential effect, $F(1,55) = 21.18$, $p < .001$, and a tendency to choose words orthographically similar to the primes, $t(55) = .09$, $p = .92$. Passive orthographic priming seems to have produced a preferential shift, but little preferential variability.

Active Priming Results

Participants took an average of 2962 ms performing the part-of-speech matching task.

The active priming results are given in Table 4. There were differences across the four types of priming, $F(3,53) = 3.74$, $p < .025$, although these did not interact with the four priming conditions, $F(9,47) = .86$, $p = .57$. The pattern of results from active priming in Experiment 1 is replicated for repetition priming: a large both-primed deficit, $t(55) = 4.49$, $p < .001$, an overall preference effect, $F(1,55) = 4.77$, $p < .05$, with a slight preference to choose the word that was **not** a repetition, $t(55) = 1.55$, $p = .06$.

Excluding repetition priming, there were no differences across the three types of orthographic priming, $F(2,54) = 2.01$, $p = .14$, and no interaction between these priming types and the four priming conditions, $F(6,50) = .40$, $p = .88$. Therefore the three types of orthographic priming are collapsed, and displayed in the lower panel of Figure 7. For the collapsed orthographic priming conditions, there was a large both-primed deficit, $t(55) = 6.89$, $p < .001$, an overall preferential effect, $F(1,55) = 21.18$, $p < .001$, but no preference for or against orthographically similar choice words (i.e. the target-primed condition was

not different from the foil-primed condition), $t(55) = .09$, $p = .92$. In other words the preference effect was due to preferential variability.

Discussion

Experiment 2 shows a general similarity between orthographic and repetition priming: with active priming both types of priming produced a both-primed deficit and either a slight preference reversal or a neutral preference; with passive priming both types of priming produced a strong preference for the related choice. This general similarity suggests that both-primed deficits and preference effects are primarily due to orthographic prime interference. This conclusion is not entirely surprising given the visual nature of the identification task. The failure to find significant differences for different degrees of semantic and phonemic similarity suggests phonemic and semantic features play a smaller role than orthographic features; these conditions produced similar results due to the overwhelming orthographic similarity.

The invariance of the passive priming results when case was changed implies that comparisons are carried out at the level of abstract orthography. The failure to obtain differences as a result of case change between primes and targets/foils is perhaps surprising conceptually (though similar results have been obtained in many other studies -- Evett & Humphreys, 1981, to name just one). It should be noted that our procedure of placing primes in different screen locations from both targets and choice words might have reduced the importance of matching visual features.

Several aspects of these results appear puzzling at first glance. In the passive priming results, the small both-primed deficit for orthographic priming as compared to the large

deficit with repetition priming suggests orthographic similarity plays less of a role in passive priming. At the same time in the passive priming results, the larger preference effect with orthographic priming suggests just the opposite. Part of the explanation may lie in the selection of words in the different conditions: In this experiment the orthographic and repetition conditions used separate pools of words. Differences in the overall perceptibility of each group of words can be inferred by comparing the neither-primed results. The participants in the active group apparently found these words equally perceptible, because neither-primed performance was the same in repetition and orthographic priming conditions, $t(55) = 1.08$, $p = .29$. However, the participants in the passive group found the perceptibility of the orthographic group of words lower than the repetition group of words, $t(51) = 3.83$, $p < .001$. We have no ready account for the existence of this passive/active by word-set interaction. Nonetheless, we shall see in the next section that using a different value of β (the target perception parameter in ROUSE) for the repetition and orthographic conditions in the passive group allows reasonably successful prediction of the results.

The ROUSE Model Applied to Experiment 2

Based upon the neither-primed analyses discussed above, a common value of β was used for the repetition and orthographic conditions in the active group, but different values of β were estimated for the repetition and orthographic conditions for the passive group. As in Experiment 1, the parameters were separately fit to the active and passive conditions since this was a between subjects variable. Table 3 gives the resultant nine parameter

estimates (4 for active priming and 5 for passive priming). The reasonably good fit of the model is illustrated by the dots in Figure 7.

For repetition priming, the both-primed deficit, the slight tendency to choose the unprimed word in the active priming group, and the tendency to choose the primed word in the passive priming group were similar to the results of Experiment 1. Thus it is not surprising that the pattern of parameter estimates is generally similar in the two experiments. In particular, in both studies the estimate of α is higher than the actual value for active priming, and lower than the actual value for passive priming. Generally speaking the estimated values of the prime similarity parameter, ρ , are reasonable: If orthography is the major determinant of interference in repetition priming and the orthographically similar words share 3/4 or 4/5 of their letters, one might expect values of ρ around .75-.80.

For the passive priming group, assuming a larger value of β for the repetition conditions than for the orthographic conditions allowed explication of a puzzle seen in the results. When the estimate of α is lower than the true value, as in the passive group, ROUSE predicts increasing separation between the target-primed and foil-primed conditions as prime similarity increases. This factor should therefore produce a greater disparity between these conditions in repetition priming (higher similarity) than in orthographic priming (lower similarity). The data show the opposite. However, ROUSE predicts the tendency to choose prime-related words will be larger for the orthographic condition because fewer target features are perceived in this condition (lower β); with fewer target-activated features, prime-activated features play a larger role, and preference is greater.

For the active priming group, the model provides a remarkably good fit. This seems surprising considering there is considerable preferential variability for both orthographic and repetition priming (i.e. large both-primed deficits), but while repetition priming produces a slight preference against a repeated choice word, orthographic priming is on average neutral with respect to the primed choice word. Naive expectations with ROUSE suggest using the same α and α' should produce preference in the same direction regardless of prime similarity. This turns out to be false; for the γ , β , and α parameters used, there is an interaction between prime similarity and the direction of preference provided $\alpha' > \alpha$ (i.e. active priming). This interaction is explained below.

ROUSE and Prime Similarity

 Insert Figures 8 and 9 about here

To take a closer look at the role of prime similarity, we produced predictions for the four critical priming conditions, for the case when the estimate of α is lower than the true value (the upper panel of Figure 8, corresponding to passive priming), and for the case when the estimate of α is higher than the true value (the lower panel of Figure 8, corresponding to active priming), for values of ρ ranging from zero to one. The parameter values are typical, in that β is set to .05 and α is set to .1 (as in all simulations γ is .02). To make the predictions clearly visible, two slightly exaggerated values of α' are used: .05 (corresponding to passive priming), and .3 (corresponding to active priming).

For passive priming, the predictions are shown in the upper panel of Figure 8. These conform to expectations. As the value of ρ drops from 1, corresponding to the switch from repetition to orthographic priming in Experiment 2 or to associative priming in Experiment 1, we see that there is no change in the neither-primed case (as must be the case since no features are shared), and that the both-primed deficit decreases (as must be the case because fewer shared features produces smaller variation in the numbers activated by the primes). The strong preference for the primed alternative with high values of similarity gradually decreases to zero as the similarity decreases to zero.

The active priming predictions, shown in the lower panel of Figure 8 do not conform to naive expectations. Predictions for neither primed and both-primed are the same as in the upper panel: Since α is the same in both panels, the neither-primed and both-primed conditions are identical; the estimate of α only affects the target-primed and foil-primed conditions. The preference predictions are the place where the failures of intuition appear: For high similarity (e.g. $\rho = 1$) there is a preference for the choice not related to the prime, as expected when $\alpha' > \alpha$. As similarity drops, however, the target-primed and foil-primed conditions change non-monotonically and the direction of preference is seen to change from a preference against to a preference in favor of prime related words.

Reference to the two factors (i.e. *Prime Activation* and *Discounting Target/Noise Activation*) discussed in association with Figure 5 helps to explain this crossover. As the similarity of the prime to a choice word decreases from 1.0, the probability of a target or noise activated feature becoming discounted gets smaller, so on many trials there will be none of this sort of discounted feature. Thus, *Discounting Target/Noise Activation* becomes unlikely as similarity decreases. As similarity drops, it also becomes less likely

that *Prime Activation* will operate: the prime is less likely to turn ON otherwise OFF features. However with the probability of target or noise activation ($\beta + \gamma - \gamma\beta = .069$) less than the probability of prime activation ($\alpha (1 - .069) = .093$), it is more likely that *Discounting Target/Noise Activation* is missing on a given trial than it is that *Prime Activation* is missing. In addition, the effect of missing *Discounting Target/Noise Activation* features is greater than missing *Prime Activation* features (i.e. for the given parameters, discounting one target/noise feature is a reduction in evidence that is approximately 3 times larger than the increase in evidence from one prime-activated feature). As *Discounting Target/Noise Activation* stochastically drops out for low ρ , a preference for prime related words becomes unavoidable.

Figure 9 is useful for highlighting the separate roles of *Prime Activation* and *Discounting Target/Noise Activation*. This figure shows log-odds distributions for three different levels of prime similarity for the case of active priming ($\alpha' = .3$). All other parameters are the same as used in the creation of Figure 8. The extremes of similarity ($\rho = 0$ and $\rho = 1$), are not shown since these can be found in Figure 6 (the same parameter values were used in the creation of Figure 6). As prime similarity increases two separate trends are observed. First, each log-odds peak is seen to spread out such that the single peak for $\rho = 0$, is replaced by a distribution of smaller peaks as ρ increases. Furthermore, the direction of this spreading is towards the preferred alternative (i.e. the target-primed peaks spread to the right and the foil-primed peaks spread to the left). This is the effect of *Prime Activation*. When otherwise OFF features are activated by the primes, this leads to subtle shifts in the log-odds towards the primed alternative. These are subtle shifts since the prime activated features are heavily discounted. Next we consider the other trend

observed in the distributions. As prime similarity increases we see the subtraction of entire central peaks and their associated spread. This is the action of *Discounting Target/Noise Activation*. Otherwise ON features are discounted due to their presence in the primes. The central peaks represent different numbers of ON features so as these ON features become discounted, peaks drop out. This dropping out preferentially affects the target-primed distribution since more features were ON in the target-primed distribution with $\rho = 0$.

In sum, Figure 9 reveals the spreading out of the distributions due to *Prime Activation* pushing the total evidence towards a positive preference, as well as the dropping out of peaks due to *Discounting Target/Noise Activation*, pushing the total evidence towards a negative preference. The effect of peaks dropping out is ultimately stronger and for high prime similarity a negative preference is obtained. Nevertheless, for low prime similarity, the effect of peaks dropping out is hardly noticeable whereas the spreading of peaks is apparent and a positive preference is obtained.

Simulations have shown that increasing β or γ or decreasing α (both of which make *Prime Activation* more likely to be missing than *Discounting Target/Noise Activation*) or increasing the number of features while keeping the other parameters constant (which makes it less likely that either factor is missing) can lessen or even eliminate the preference crossover. Note that this does not imply there is something special about the number of features we've chosen; even with a large number of features, the crossover can be reinstated if the activation parameters are sufficiently reduced. However, the presence of a crossover in the empirical data is indicative that visual noise (γ) is necessarily low. If γ is increased, then it becomes all but impossible to produce the crossover in the model.

A post-hoc analysis of the data in Experiment 2 revealed the qualitative trends seen in the simulations of Figure 8. Even though we have modeled the neither-primed condition as if no features overlapped, in fact the primes had a random relation to the choice words, so that there were occasional matches of one or more letters in the same position within a letter string. Thus from the neither-primed trials we extracted trials on which 1, 2, or 3 out of 5 of the letters randomly matched and were in the same position as letters in the target. These cases were combined and called ‘low orthographic similarity’ (contrasted with the high orthographic similarity conditions of Experiment 2 in which all but one letter matched). For passive priming, there was a preference for the related choice word with ‘low orthographic similarity’ and in keeping with Figure 8, this preference was less than that seen for high orthographic similarity and repetition priming. However, for active priming, in keeping with Figure 8, there was a preference in favor of the related choice word for low orthographic similarity. Thus there is considerable evidence from this study supporting the various predictions illustrated in Figure 8. We realize, however, that a post-hoc analysis is not very conclusive; therefore in Experiment 6 we manipulate prime similarity directly (and, as predicted, we observe a preference crossover with active priming.).

Experiment 3: Repetition Priming and Choice Word Similarity

Ratcliff and McKoon (1997) have used a paradigm similar to ours to explore long-term repetition priming. In their long-term priming paradigm, words were studied in a first phase; in the second phase of the experiment, test words were briefly flashed followed by a

2-AFC test. They varied the orthographic similarity of the two choice words and obtained results that tightly constrained possible models. We carry out such a manipulation in the present experiment. Before turning to our study, it is useful to review the Ratcliff and McKoon findings.

Typical long-term priming results with other testing procedures (e.g. lexical decision or naming) show facilitation with repetition priming, but not with associative or other sorts of priming (e.g. Ratcliff, Hockley, & McKoon, 1985; however for an alternate account see Becker, Moscovitch, Behrmann, & Joordens, 1997). For repetition priming, changes in the appearance of the choice word relative to the study word reduce the magnitude of the priming effect, but do not eliminate it (e.g. Bowers & Michita, 1998).

In most long-term repetition priming studies the facilitation observed could be due to a type of preference to choose a recently encountered word. This led Ratcliff and McKoon to use the 2-AFC design, to determine if facilitation is due to a perceptual benefit or preference. They found no change in the both-primed condition, and the average of the target-primed and foil-primed conditions was about equal to performance in the neither-primed condition. Both results suggest there was no perceptual facilitation for the target flash. However, there was a preference for repetitions since the target-primed condition was higher than the foil-primed condition. The fact that this advantage occurred when the choice words were orthographically similar, and was greatly reduced when the choice words had differing orthography, was a key factor in leading Ratcliff and McKoon to conclude that long-term repetition priming was a matter of "bias" rather than improved perception⁵. Ratcliff and McKoon (1997) developed a model to account for this pattern of

results, and Schooler, Shiffrin, and Raaijmakers (1997) developed an alternative model to do the same, based on Bayesian principles similar to those used in this thesis.

Because the pattern of results observed by Ratcliff and McKoon differed markedly for similar and dissimilar choice words, and because these differences were critical to constraining possible models of the results, we decided to test the ROUSE model by varying the similarity of the choice words in our present short-term priming paradigm. Experiment 3 uses only repetition priming. It uses the word pairs of Experiment 2, but includes conditions in which the orthographically similar word pairs of that study appear as choice words. Instead of using pairs of dissimilar words selected to be as dissimilar as possible, as was done in Ratcliff and McKoon's study, we created our dissimilar word pairs by randomly pairing two words. We term this neutral similarity.

Method

There were 55 participants in the passive priming condition and 56 in the active priming condition.

Repetition priming was used throughout. The four categories of word pairs found in Experiment 2 were reused; in the 3 sets of orthographically similar choice word conditions, the members of a given pair were presented as choice words (one being the target, the other the foil). The neutral similarity choice word conditions were identical to the repetition conditions found in Experiment 2. All other procedures were as in Experiment 2. The active priming group used the same part-of-speech matching task. In Experiment 2 the letter case manipulation was only administered to the passive priming group. In order to test the generality of case indifference, Experiment 3 employs the letter

case manipulation for the active priming group; as in Experiment 2, this condition combines the *orthographic* and *orthographic and phonemic* word pairs and randomly assigns word pairs to conditions in which primes appear in lower case or conditions in which primes appear in upper case (resulting in same case and case switch between prime and target). For the passive priming group the word pairs were kept in their original categories and all words appeared in lower case. The conditions are illustrated with examples using particular words in Table 5.

Passive Priming Results

 Insert Table 5 about here

The results are given in Table 5. There were three types of orthographic choice word similarity; these did not differ statistically, $F(2,53) = 3.11$, $p = .05$, and interacted only weakly with priming condition, $F(6,49) = 2.84$, $p < .025$. Inclusion of the neutral choice word similarity conditions led to differences, $F(3,52) = 16.59$, $p < .001$, and these differences interacted with the four priming conditions, $F(9,46) = 3.05$, $p < .01$. This pattern was largely due to better performance for the neutral than orthographically similar choice-word conditions (collapsing across priming conditions, $t(55) = 10.42$, $p < .001$). This is natural because the decision is more difficult when the choice words are more similar. Because the differences among types of orthographic similarity were small, these conditions were combined, and the results graphed in Figure 10.

Insert Figure 10 about here

There was a highly significant both-primed deficit in the orthographically similar choice-word condition, $t(54) = 6.04$, $p < .001$. However, in the neutral similarity conditions, the small both-primed deficit did not reach significance, $t(54) = .9$, $p = .37$. There were preferential effects for both the orthographically similar, $F(1,54) = 40.44$, $p < .001$, and neutral similarity, $F(1,54) = 111.57$, $p < .001$, conditions: Participants tended to choose the repeated word for both types of choice word similarity (neutral similarity: $t(54) = 6.75$, $p < .001$; orthographically similar: $t(54) = 8.24$, $p < .001$). These results largely replicate those from the passive groups in Experiments 1 and 2.

Active Priming Results

Participants in the active priming group took an average of 2811 ms to perform the part-of-speech matching task.

Three types of orthographic priming were used: two types with or without a case change and a third type using the *orthographic and semantic* word pairs. The results are shown in Table 5. There were differences across the three types of orthographic choice-word similarity, $F(2,54) = 5.18$, $p < .01$, but these differences did not interact with priming condition, $F(6,50) = 2.11$, $p = .07$. In all cases the qualitative trends across priming conditions were the same. Presenting primes in a different case than the target and choice words produced differences that interacted weakly with the four priming conditions, $F(3,53) = 2.93$, $p < .05$, but as seen in Table 5, this was due to quantitative, not qualitative

differences; these results again support the conclusion that feature comparisons occur largely at the level of abstract orthography. In light of the foregoing analyses, the results are collapsed across the three types of orthographic similarity and the case change manipulation, and graphed in Figure 10.

As with the passive group, there were differences with choice word similarity when the neutral choice word similarity conditions were included, $F(3,53) = 37.21$, $p < .001$, and these differences interacted with the four priming conditions, $F(9,47) = 4.86$, $p < .001$; these differences were due to lower performance with orthographically similar choice words, the expected result when the choice words are more similar.

For both neutral similarity, $t(55) = 2.92$, $p < .005$, and orthographically similar choice words, $t(55) = 7.41$, $p < .001$, there were deficits in the both-primed conditions. Likewise, there were preferential effects for both neutral similarity, $F(1,55) = 4.08$, $p < .05$, and similar choice words, $F(1,55) = 79.54$, $p < .001$. Comparing the target-primed and foil-primed conditions, participants preferred the repeated word when the choice words were orthographically similar $t(55) = 6.12$, $p < .001$ (it is important to note that this result is opposite to that found in the active group in Experiments 1 and 2); there was no difference between these conditions when the choice words were of neutral similarity, $t(55) = .05$, $p = .96$ (in keeping with the preference removal or reversals seen in Experiments 1 and 2 with active priming).

Discussion: ROUSE and Choice Word Similarity

For active priming, the observation of a preference for repeated words only for similar choice words and not for dissimilar words is analogous to the Ratcliff and McKoon (1997)

results, but does not seem at first glance to be in accord with the predictions of ROUSE (or the results from the first two experiments). As we demonstrate next, however, the model provides a surprisingly good account of the findings.

How can ROUSE predict a preference for repeated words when the choices are similar, even though α' is larger than α ? In the discussion associated with Figure 5, the effect of priming was separated into two components: *Prime Activation* was shown to promote a preference for related words regardless of α' (decreased in magnitude for higher α' , but still in favor of related words), whereas *Discounting Target/Noise Activation* was shown to reduce the evidence in favor of the primed alternative allowing the removal of or even a reversal in the direction of preference.

A closer examination of these two factors reveals that *Discounting Target/Noise Activation* crucially depends upon the existence of noise-activated features. Consider the situation without noise ($\gamma=0$). For the target-primed condition, foil features can only be activated by noise; without noise all foil features are OFF. In this no noise situation, discounting the ON target features results in a reduction of target evidence but does not affect $P(C)$ since even discounted features provide more evidence than the entirely OFF foil features; in this situation, $P(C)$ will be determined by whether any target features are ON since even a single ON feature, even if discounted, will result in choosing the target. In other words *Discounting Target/Noise Activation* only affects performance when there are ON features in the unprimed word. For the foil-primed condition, the removal of noise-activated features also disables the effect of *Discounting Target/Noise Activation*. Without noise, foil features are only activated by the primes. So the effect of *Discounting Target/Noise Activation* has been eliminated because there are no features of the foil that

are activated by the target and/or noise. For both the target-primed and foil-primed conditions the effect of *Prime Activation* is preserved in the no noise situation; additional features in the primed alternative are activated, receiving a discounted level of evidence, and this results in a preference for primed words.

To see how this dependence on noise-activated features results in a preference for repeated words when the choice words are similar, recall that features shared by both alternatives are non-diagnostic and drop out of the decision process. Experiment 2 demonstrated that orthographic features dominate the decision process so it can be expected that only a small number of features are diagnostic with orthographically similar choice words. With only a small number of relevant features, performance will decrease (i.e. there are fewer flash activated features) but also noise will decrease (i.e. there are fewer noise-activated features). Considering the low value of γ used in the simulations, on many trials there will be no noise-activated features. For these trials, provided there are some prime-activated features (which will most likely be true since $\alpha \gg \gamma$), a preference for related words will result regardless of the level of discounting.

To sum up, the effect of increasing choice word similarity is to reduce the number of relevant features. This makes it unlikely that any features are turned on by noise which in turn disables the preference removal/reversal due to *Discounting Target/Noise Activation*.⁶ The dependence of preference removal/reversal on the existence of noise-activated features has been demonstrated with simulations. If the noise parameter is set too high, then the preference change with similar choice words is eliminated; if $\alpha' > \alpha$ and $\gamma \gg 0$, a preference against prime related words occurs regardless of choice word similarity. As with the prime similarity crossover discussed in Experiment 2, the

observation of a change to a preference in favor of repeated words (under active priming) with highly similar choice words is suggestive of a low visual noise situation.

The ROUSE Model Applied to Experiment 3

The predictions of the ROUSE model, based on the best fitting parameters found in Table 3, are given in Figure 10. The model's predictions mimic the qualitative patterns of results, and come close to many of the quantitative observations. The crucial prediction of the model occurs for the active priming group: The model predicts both 1) for randomly similar choice words, no differential preference for the repeated choice word, and 2) for highly similar choice words, a preference for the repeated choice word. It is particularly striking that this second prediction differs from those for the active groups in Experiments 1 and 2, even though in all three cases α' is estimated to be higher than α . That is, the estimate of prime activation is higher than the actual value in all three studies, yet the direction of preference is predicted to reverse in this study, due to the similarity of the choice words. Even more compelling is that the reversal is observed within subject and the model predicts the results using the same α and α' values for both directions of preference. That this unanticipated prediction matches the results lends some credence to the model.

It is important to note that the orthographic similarity parameter, ρ , performs a different function in Experiment 3 than it did in Experiments 1 and 2. In Experiments 1 and 2, ρ determined the similarity of prime to choice word, and was set at a level sufficient to produce the observed relation between repetition and associative priming or between repetition and orthographic priming. In Experiment 3, ρ determines similarity between the

two choice words and its value is most important for determining the proportion of features in the decision process (which determines performance levels and the switch in the active priming preference results). Since the prime to choice word similarity and choice word to choice word similarity could depend on attentional factors in different ways, the value of p might differ between studies, even when letter overlap is identical in the two cases.

As seen in Figure 10, in a few instances the fit to data is quantitatively awry. The most important discrepancy involves the deficit for both-primed: The data show a larger deficit for the orthographic conditions, whereas the model predicts a slightly larger deficit for the neutral conditions. Because the both-primed condition in this study presents two primes that differ in one letter position, subjects may attend more to that position in the primes, raising the value of α for that feature, the only diagnostic feature. A higher value of α would increase the variability of priming, causing performance to decrease.

A different but equally plausible account of the pattern of both-primed deficits involves subject variability. Despite our efforts to fix target duration at a level that would equate the performance levels across participants, there was in practice a wide range of performance levels. For instance, of the 56 participants in the active priming group, 4 participants had an average performance above 90% while 11 others had an average performance less than 60%. Such variability can alter the predictions in a number of ways. In particular, subject variability can especially decrease predicted both-primed performance when the choices are similar: The poorly performing participants remain poor when the task difficulty is increased by imposing similar choices; the better performing participants are hurt when similarity increases making them more sensitive to preferential

variability. This informal reasoning was verified through additional simulations. We introduced a positive correlation among the target-activated features by randomly sampling on each trial the target flash activation parameter, β , from a distribution. This approach improved the fit shown in Figure 10. Lacking appropriate data to constrain the form of the distribution for β , we decided not to report these predictions in detail, and leave such modifications of ROUSE for future work.

Experiment 4: Perceptual Benefits of Associative Priming

The experiments thus far have not demonstrated perceptual benefits due to priming, and the results have been predicted rather well by a model based solely on preferential factors. For example, Experiment 1, using strong association strengths for associative priming, failed to find differences between the neither-primed and both-primed conditions, although repetition priming produced large both-primed deficits. This null effect in the associative case could have resulted from a trade off of preferential variability with perceptual facilitation, both induced by priming. Nevertheless the pattern of data, and the success of ROUSE in fitting both associative and repetition priming data without assuming any perceptual facilitation suggest that any perceptual facilitation due specifically to associative priming was fairly weak.

There are a number of reasons, however, to suspect that this is not the entire story. For example, Masson and Borowski (1998) obtained results that clearly indicated the existence of perceptual facilitation in their associative priming study. Thus we decided to explore more carefully the possibility of facilitative effects for associative priming. In

particular, we decided to switch to a single prime word, on the theory that the use of two prime words might dilute or disrupt associative priming effects. Experiment 4 therefore uses a single prime and tests associative priming only. In the control condition, neither choice word is related to the prime. In the both-primed condition, both choice words are related and have the same number of letters. For instance, the prime MISTAKE is followed by the target WRONG, followed by a 2-AFC between the choice words WRONG and ERROR. It is more difficult to find primes and strong associates that meet these criteria, yet the lack of a competing prime word might more than offset this cost.

Also in an attempt to increase the efficacy of priming, we decided to place the primes in the same screen position as the subsequent target flash. Of course such a procedure introduces the possibility of forward and backward masking between prime and target. In an unpublished series of experiments we found that forward masks (e.g. primes in the same location as the target), initially gain effectiveness with increased duration out to around 100ms, but then lose effectiveness as duration is further increased to 2000ms. This u-shaped function of forward mask duration on accuracy could be explained in terms of increasing stimulus energy for shorter durations and the inverse duration effect (Coltheart, 1980; Hogben & Di Lollo, 1974) for longer durations. Whatever the explanation, the findings suggest it would be useful to vary prime durations, and we did so in Experiment 4, even including two conditions in which the prime follows the target. Because we expect performance to vary due to masking in the present study, we also include control conditions in which a pattern mask is presented instead of a prime word. A final control condition presents no prime, the display sequence in this condition consisting of the target flash followed by a backward pattern mask.

The target-primed and foil-primed conditions are not used in this study. Only the passive priming procedure is adopted.

Method

94 Indiana University undergraduates participated.

264 prime/associate1/associate2 (e.g. MISTAKE/WRONG/ERROR) triples were created. Prime and associates ranged from 3 to 7 letters in length with prime length and associate length often differing. Both associates of a prime contained the same number of letters. Since each associate could appear as target or foil, association strength was counterbalanced between participants. The average association strength was .145. Unlike all other experiments, the screen refresh rate was 70Hz (14.3ms) and 3 different target durations were used in lieu of setting target duration separately for each participant. 12 point Geneva font was used for all displays. Since all words appeared in upper case and targets could appear following or prior to primes, participants had no way of identifying which word was the target until the 2-AFC.

There were 66 conditions in the experiment. 63 of these resulting from 3 target durations by 3 types of primes by 7 prime durations. In addition there were 3 target durations when *no prime* was presented. Each of the conditions occurred 4 times during the experiment. Target durations were 29, 43, and 57ms. The three types of primes were pattern mask, neither-primed (created by re-pairing prime and associates), and both-primed (intact prime-associate presentation). Target and foil were always the two associates of a single prime (even on neither-primed and pattern mask trials). In other words participants always chose between two indirectly related words. The 7 prime

durations consisted of 5 conditions in which the prime/mask preceded the target and 2 conditions in which the prime/mask followed the target. The 5 preceding times were 29, 57, 114, 457, and 1829ms. The 2 following times were 29 or 57 ms.

Presentation order was 500ms fixation, 500ms blank screen, prime (if the prime preceded), target, prime (if the prime followed), 500ms pattern mask, 500ms blank screen, and the 2-AFC screen until response. Except for the 2-AFC, all items appeared in the same location in the center of the screen. Unlike previous experiments, the 2-AFC consisted of one word above and the other below the central position. It is important to note that in the *no prime* condition and in the *prime following target* conditions, there was no forward mask on the target. All conditions backward masked the target with either a pattern mask or word. The final 500ms pattern mask was not reduced by the target duration as in other experiments. If the prime appeared following the target, the duration of the final pattern mask was reduced by the duration of the prime. This served to equate the time between offset of the target and presentation of 2-AFC.

Two experimental blocks were preceded by 8 practice trials. The 8 practice trials were created from 8 additional associative triples and conditions were selected to accustom participants to the types of presentations possible within the experiment. Since fixed target durations were used, there was no block of trials for threshold determination. Participants viewed the primes under passive priming conditions. However, since it could not be known in advance whether the target was the first or second word within the presentation sequence, primes were probably treated as potential targets.

Results and Discussion

There was a main effect of target duration, $F(2,92) = 320.17$, $p < .001$, but no interactions with priming condition, $F(2,92) = 1.66$, $p = .20$, or with prime duration, $F(12,82) = 1.29$, $p = .24$, so the results are collapsed over target duration. Figure 11 shows the resultant accuracy results for the both-primed and neither-primed conditions. Analyzing the neither-primed and both-primed priming conditions (i.e. excluding the pattern mask conditions), there were main effects of prime duration, $F(6,88) = 33.29$, $p < .001$, and prime condition, $F(1,93) = 22.78$, $p < .001$. There was no interaction between priming condition and prime duration, $F(12,82) = 1.16$, $p = .33$ (i.e. the difference between neither-primed and both-primed was in the same direction for all prime durations).

 Insert Figures 11 and 12 about here

The main effect of prime condition (i.e. both-primed better than neither-primed) was observed at all prime durations (although this difference was not significant at every duration). Similar to the Masson and Borowsky (1998) study with associative priming, this result suggests a gain in perceptual processing of the target (according to our definition of perceptual facilitation). This result was stronger than that achieved in Experiment 1, despite weaker association strengths in the present case. Most likely the failure to find a both-primed benefit in Experiment 1 in light of our success in finding a both-primed benefit in this experiment is due to the use of two versus one prime. (It is also possible that the use of the same location for prime and target could have increased the

size of effect, but see Experiment 6 in which the both-primed benefit is replicated with primes in different locations).

In any case, this improvement in performance cannot be handled by the present version of the ROUSE model. To account for these results a mechanism for perceptual facilitation is required. Potential mechanisms are discussed in the section within the General Discussion titled, '*Short-term Priming Theories and Perceptual Effects*'. It is important to note that the magnitude of the perceptual improvement is small relative to many of the preferential effects we have observed; the improvement is only around 3%, whereas various preferential effects exceeded 10% in the first three experiments. Because most word priming experiments do not include anything analogous to the both-primed (i.e. preference controlled) condition, it is likely that the greater part of the observed facilitation in those studies corresponds to the target-primed enhancement in the passive conditions of Experiments 1-3. Hence it seems likely that the priming benefits in most previous associative priming experiments are due largely to preferential rather than perceptual effects.

Perhaps surprisingly, there was a both-primed improvement for backwards priming at the 57ms prime duration, $t(93) = 2.42$, $p < .01$. In choosing the associates for this experiment, only the forward direction (i.e. given the prime, produce the associate) association strengths were considered; the backward association strengths were unconstrained. Even if the backward association strengths were large, it is surprising that perceptual facilitation for a target word should occur as a result of a word presented after target presentation. This suggests a fair amount of temporal overlap in the visual processing of words.

As measured by separate t-tests, the both-primed advantage is significant at the .05 level or better for all prime durations except the longest prime preceding target duration, 1829 ms, and the shortest prime following target duration, 29 ms. The trend is one of initially increasing improvement with increasing prime duration and then decreasing improvement for longer durations. Since the facilitation disappears at long prime durations, perceptual facilitation with associative priming appears to be a short-term priming phenomenon.

The significant facilitation at 29 ms of prime preceding target demonstrates perceptual facilitation for what could be termed subliminal priming; this is a short enough duration that on many trials (but not all), participants will be unaware of the prime's identity. In the General Discussion, the predictions of the ROUSE model for subliminal priming are discussed in more detail.

Next the pattern mask conditions are considered. Figure 12 shows the accuracy results for the pattern mask and neither-primed conditions. The neither-primed conditions in Figure 12 are identical to those seen in Figure 11 and are re-plotted for the sake of comparison. Analyzing these two sets of conditions, there was a main effect of target duration, $F(2,92) = 293.8$, $p < .001$, that interacted with prime condition (pattern mask vs. neutral word), $F(2,92) = 4.94$, $p < .01$, and with prime duration, $F(12,82) = 2.16$, $p < .025$. Nevertheless, the qualitative trends were the same regardless of target duration as revealed by separate plots and t-tests and therefore target duration is collapsed in Figure 12.

There was a main effect of prime condition, $F(1,93) = 10.73$, $p < .0025$, and a main effect of prime duration, $F(6,88) = 31.72$, $p < .001$, and these two variables strongly interacted, $F(6,88) = 23.32$, $p < .001$. The neither-primed condition was significantly better

than the pattern mask condition when the prime/mask preceded the target by 29 or 57ms, but the two conditions were not different for all other preceding durations. When the prime/mask followed the target, the pattern mask condition significantly outperformed the neither-primed condition for both durations. This pattern of results held true when each of the three target durations were considered separately.

The results for prime/mask following target are considered. Pattern mask better than neither-primed is not surprising in light of the normally observed backward masking characteristics: It has been known for some time that letters and words are more effective backward masks than pattern masks, (e.g. Taylor & Chabot, 1978). Since the 29 or 57ms pattern mask is subsequently replaced by a highly similar long duration (500ms) pattern mask, there is no effect of mask duration. In either case the sequence of events is target flash followed by a long period of pattern masking. The decrease in performance as a function of neither-primed word duration is also explained in terms of mask type. Since words are better backward masks than pattern masks, performance drops as word duration increases: Longer duration word masks have more stimulus energy and therefore are more effective backward masks.

The prime/mask preceding target results are puzzling although in keeping with unpublished experiments we have performed looking at the masking characteristics in perceptual identification. For both the neither-primed and pattern mask conditions there is a u-shaped pattern such that as preceding duration increases, performance initially decreases and later, for very long durations, performance is seen to rise. Since this pattern holds for both conditions, lexical level effects can be ruled out as a potential explanation: For the most part, the nonmonotonicity must be due to visual masking. A coherent

explanation can be given in terms of visual persistence. It is known that for short durations (<200ms) visual persistence increases with increasing duration whereas for long durations (>1000ms) visual persistence decreases with increasing duration (Coltheart, 1980; Hogben & Di Lollo, 1974). Assuming forward masking occurs due to the simultaneous activation of the forward mask and subsequent target, this same u-shaped pattern is expected as a function of forward mask duration.

The only remaining mystery is neither-primed better than pattern mask for 29 and 57ms of preceding duration. If words are better masks than pattern masks for backward masking, the same might be expected for a forward mask. Instead the opposite is observed. One possibility is that the neutral word forward mask initiates attentional process since any word is a potential target. There is evidence that it takes around 100ms to open an “attentional spotlight” (e.g. Sperling & Weichselgartner, 1995) so the short duration forward word mask might serve to turn on such a spotlight in a timely manner. Whatever the explanation, further work is required to fully explicate these complex masking and attentional processes.

In sum, use of a single associative prime can give rise to a small, but reliable, both-primed benefit. Perceptual factors, not presently included in the ROUSE model, are necessary to predict such a result.

Experiment 5: Associative Priming and Choice Word Similarity

In Experiment 3, orthographically similar choice words produced a larger both-primed deficit than neutrally related choice words (although the simplest version of ROUSE predicted the reverse). Experiment 5 tests the effect of orthographically similar choice words upon the both-primed benefit found in associative priming. In this study a single prime is presented for 250ms under passive priming conditions; conditions that maximized the both-primed benefit in Experiment 4. In the both-primed condition the single prime must associate to two choice words, tending to produce semantic similarity between the choice words. Nonetheless we decided to vary the degree of both semantic and orthographic similarity of the choices in Experiment 5. The target-primed and foil-primed conditions are included in this study in order to assess any relationship between preference and the expected both-primed benefit.

Method

69 Indiana University undergraduates participated. In the associative norms (Nelson, McEvoy, & Schreiber, 1994) variants of the same word are collapsed. For instance if participants are given the word RESERVE, some will respond HOLD and others might respond HELD (i.e. HOLD and HELD are collapsed into a single response category). Experiment 5 takes advantage of this in order to create a both-primed associative condition in which the choice words are orthographically as well as semantically similar. Mostly through tense changes such as HOLD/HELD, 105 prime associate triples were created. Primes were 3 to 10 letters in length and associates were 4 to 9 letters in length. The two variants of the same word differed in only one letter. On every trial, regardless of condition, one variant of the same word associate (e.g. HOLD or HELD) was fixed as the

target word. Depending upon which condition was being tested, the appropriate prime and foil words were then selected. Unlike Experiment 3, fixed target and foil words were used to create orthographically and/or semantically dissimilar choice words (in Experiment 3 these were created by randomly pairing words). For this reason we use the term dissimilar choice word similarity for this experiment.

There were 15 conditions in the experiment: Seven of them were: both-primed with dissimilar choice words; both-primed with semantically similar choice words; both-primed with orthographically and semantically similar choice words; target-primed with dissimilar choice words; target-primed with orthographically similar choice words; foil-primed with dissimilar choice words; and foil-primed with orthographically similar choice words. Each of these seven conditions had a neither-primed baseline, also using the same pool of words, making 14 conditions. Five additional words per triple were required to fill these conditions. (Suppose for example in the following discussion that the triple consisted of the prime HUM and the target SONG with its variant SING). Three of these five additional words were a semantically similar associate foil word (TUNE), a dissimilar foil word (HOUR), and an orthographically similar (also one letter different) foil word (LONG). In order to create a both-primed test with orthographically and semantically dissimilar words, it was necessary to use a different prime word. For instance the prime THEME associates to the target SONG as well as to the foil IDEA. The fifteenth condition was a thus a re-test of the target primed condition with dissimilar choice words using this new prime word (prime: THEME, target: SONG, and foil: HOUR). This was done to measure whether the two pools of prime words produced equivalent amounts of priming/preference. Table 6 lists the basic 15 conditions illustrated with the specific

SONG/SING example. The average association strengths were as follows: targets = .16, semantically similar foils = .12, targets when primed with the alternate prime used for testing neutral similarity choice words = .11, and dissimilar foils = .11.

Each of the 15 conditions was repeated 7 times scattered within the block of experimental trials. Primes appeared for 250ms. The 2-AFC screen consisted of choice alternatives presented above and below each other. All other aspects of the experiment were similar to the passive version of Experiments 1.

Results and Discussion

The mean accuracy results are found in Table 6. As hoped, the two pools of prime words were equally effective; there was no difference in performance between the target-primed condition with dissimilar choice words when the primary primes (.803) and alternate primes (.795) were used, $t(68) = .34$, $p = .74$. Since a different pool of foil words was used for each of the primed conditions, separate ANOVA's were run for the both-primed, target-primed, and foil-primed sets of conditions.

 Insert Table 6 about here

For the both-primed conditions, there was no main effect of priming condition (i.e. both-primed versus neither-primed), $F(1,68) = 1.41$, $p = .24$, no main effect of choice word similarity, $F(2,67) = 1.39$, $p = .26$, and no interaction between priming condition and choice word similarity, $F(2,67) = .04$, $p = .97$. The failure to obtain a benefit for the both-primed conditions is surprising, since these conditions are in many respects a replication of

Experiment 4. Two possibilities come to mind. Although the association strengths in this experiment were comparable to those found in Experiment 4, it is not clear that this applies equally to both variants of the same word (e.g. SING versus SONG). The associative norms do not include a breakdown by word variant, so one variant might be much lower in association strength than the other. This would add additional preferential noise into the decision process offsetting any perceptual facilitation. The second, perhaps more likely, explanation is insufficient power to observe the small improvement in performance due to associative priming: Even though the results did not reach significance, for all three types of choice word similarity, the both-primed conditions exceeded the corresponding neither-primed conditions. This experiment included 483 data points per condition whereas Experiment 4 included 1128 data points per condition (after collapsing target duration).

Within the foil-primed conditions, there was a main effect of priming condition (neither-primed versus foil-primed), $F(1,68) = 8.52$, $p < .01$, and a main effect of choice word similarity, $F(1,68) = 8.03$, $p < .01$, but no interaction between priming condition and choice word similarity, $F(1,68) = 1.00$, $p = .32$. Within the target-primed conditions, there was a main effect of priming condition (neither-primed versus target-primed), $F(1,68) = 13.82$, $p < .001$, and a main effect of choice word similarity, $F(1,68) = 6.38$, $p < .025$, but no interaction between priming condition and choice word similarity, $F(1,68) = .003$, $p = .95$. Since the foil-primed conditions decreased as compared to their neither-primed conditions while the target-primed conditions increased as compared to their neither-primed conditions, there was a tendency to choose the word related to the prime (i.e. a

positive preference). Furthermore, this tendency was relatively unaffected by choice word similarity.

In summary, the results of Experiment 5 emphasize the small magnitude, and difficulty of obtaining, a both-primed benefit with associative priming. In both this (passive priming) study and the passive group in Experiment 3 there was a tendency to choose the related word regardless of choice word similarity.

Experiment 6: Associative and Orthographic Priming

Experiment 4 succeeded while Experiments 1 and 5 failed to obtain a both-primed benefit with associative priming. Experiment 6 tests whether finding such a benefit depends on visual location and association directionality between prime and target, and upon active or passive prime processing. In addition, this study explores the presence and direction of any preference for prime related words as a function of the type of prime (i.e. associative versus orthographic priming); in one set of conditions a single prime word is used that is orthographically similar to both choice words.

To explore the effect of visual location, primes appear above and below the central location (as in Experiments 1-3) but the same prime word appears in both locations.

Association directionality is broken into 3 different categories. In a *symmetrical* association, the prime associates to the target and/or foil that associates back to the prime. In an asymmetrical *forward* association, the prime associates to the target and/or foil that does not associate back to the prime. In an asymmetrical *backward* association, the prime associates to neither choice word yet the target and/or foil associates back to the prime.

Most experiments select words based upon the forward association strengths without reference to the backward association strengths (however see Thompson-Schill, Kurtz, & Gabrieli, 1998). Therefore traditional priming studies use some mixture of the *symmetrical* and *forward* categories. Different theories of representation will predict different results for these priming categories. For instance in ROUSE all that matters is shared semantic features. Association directionality will play a role only in as much as it reflects differential degrees of semantic similarity. ROUSE predicts that these three categories of priming should produce qualitatively similar results.

Method

There were 62 participants in the active priming group and 87 in the passive priming group.

The procedures in Experiment 6 were similar to those of Experiments 1, but only one word appeared as the prime instead of two. The one prime word was repeated in two locations on the screen corresponding to the two locations where separate prime words were displayed in other experiments. Chin rests were not used.

Four categories of 40 triples were created: three categories with a prime and two associates and a fourth category with a prime and two orthographically similar words. For the associative categories, all associates were 4 or 5 letters in length and primes were 3-5 letters in length. In the *forward* category, primes associated to two associates with an average strength of .168. Searching through the norms with these associates as potential primes revealed that they did not associate back to the primes with any known strength. This does not necessarily imply that they were semantically dissimilar. For example the

prime ASHES associates to DUST, but DUST does not associate back to ASHES. In selecting these asymmetrical associates it was required that all associates exist in the norms as primes. In the *backward* category two choice words associated to the prime, but the prime did not associate to either choice word. The average association strength was .260. In the symmetrical category, association strengths were found in both directions with an average forward strength of .217 and a backward strength of at least .07. The orthographic category was created from 5 letter words. The prime was required to share 4 out of 5 letters (in the same positions) with each of the choice words (but not necessarily the same 4 letters). An additional pool of 336 5-letter words was created for use in the practice sessions and threshold duration block.

A necessity created by the use of a single prime design was differential choice word similarity for the neither-primed and both-primed conditions as compared to the target-primed and foil-primed conditions. For instance, the prime CURVE would be followed by a choice between CARVE and CURSE in the orthographic, both primed condition. In this example it is impossible to test the target-primed condition with this same high level of choice word similarity; the foil would necessarily be similar to the prime. To create the target-primed and foil-primed conditions neutral choice words were randomly selected. Despite this difference, the neither-primed and both-primed conditions can be compared to each other, and likewise the target-primed and foil-primed conditions can be compared to each other. To a lesser degree the analogous situation exists for associative priming (i.e. the choice words are necessarily somewhat semantically similar to each other only in the both-primed and neither-primed conditions).

With only one prime, a new task was necessary for active priming. Participants performed a simple affect task by rating each prime word as positive or negative. For passive priming, the prime was presented for 250 ms (see Experiment 4 for the rationale behind this duration). Each of the four categories of words was tested in each of the 4 basic priming conditions. These 16 conditions each appeared 10 times distributed throughout the experiment. Table 7 contains specific examples of these 16 conditions. All other procedures were as explained in the General Method section.

Results

For the active priming group, participants took an average of 1179 ms performing the affect rating task.

There were no differences across the three types of associative priming: for the passive group, $F(2,85) = 1.01$, $p = .37$, and for the active group, $F(2,60) = .55$, $p = .58$. There were also no interactions between these priming types and the four priming conditions: for the passive group, $F(6,81) = 1.21$, $p = .31$, as well as for the active group, $F(6,56) = .71$, $p = .65$. Therefore the three types of associative priming are collapsed.

 Insert Table 7 and Figure 13 about here

Figure 13 shows the priming results separately for orthographic priming and the average of the three types of associative priming and Table 7 shows the separate results. Including orthographic priming, there were differences across the four types of priming: for the passive group, $F(3,84) = 27.78$, $p < .001$, and for the active group, $F(3,59) =$

11.86, $p < .001$. These differences interacted with the four priming conditions: for the passive group, $F(9,78) = 6.40$, $p < .001$, and for the active group, $F(9,53) = 2.99$, $p < .01$. Since there were no differences between the three types of associative priming, these differences with the inclusion of orthographic priming reflect differences between associative and orthographic priming.

The results for the orthographic priming conditions appear similar to the repetition priming results of Experiments 1-3, for both the passive and active groups. A both-primed deficit was observed for the passive group, $t(86) = 6.74$, $p < .001$, as well as for the active group, $t(61) = 2.27$, $p < .025$. The usual check for the existence of preferential effects cannot be performed since the neither-primed / foil-primed and target-primed / both-primed comparisons are confounded with choice word similarity. However the tendency to respond with a word related, or unrelated, to the prime is also indicative of preferential effects, and these were clear: for the passive group there was a tendency to choose the word orthographically related to the prime, $t(86) = 3.63$, $p < .001$, whereas for the active group there was a tendency to choose the word not orthographically related to the prime, $t(61) = 2.56$, $p < .01$.

For associative priming in the passive group there was a both-primed benefit (although small in magnitude), $t(86) = 2.78$, $p < .005$; for the active group, there was no both-primed deficit or benefit, $t(61) = .48$, $p = .64$. There was a tendency to choose the word associatively related to the prime regardless of the manner in which the prime was viewed: for the passive group, $t(86) = 7.97$, $p < .001$, and for the active group, $t(61) = 2.86$; $p < .005$.

Discussion

The both-primed associative benefit found in Experiment 4 is replicated in this study with the prime(s) displayed in a different location than the target: Associative priming with a single passively viewed prime presented in two locations for 250ms produces a small both-primed benefit. This benefit with passive priming as well as the tendency to choose prime related choice words was observed regardless of the direction of association (see Table 7). The uniformity of priming regardless of association directionality suggests a theory, such as ROUSE, in which feature overlap, not production, is the crucial determinant of preferential and perceptual aspects of priming. This result is in agreement with a recent study finding associative priming in lexical decision regardless of association direction (Thompson-Schill, Kurtz & Gabrieli, 1998). In our experiment, the perceptual benefit occurred across different visual locations, so the finding of such a benefit is not dependent upon prime and target occupying the same screen location. As we have noted before, the ROUSE theory in its present form cannot predict both-primed benefits. Possible mechanisms for such perceptual benefits are taken up in the General Discussion.

Similar to the repetition priming results and most of the orthographic priming results in Experiments 1-3, orthographic priming in this study produced both-primed deficits for both the passive and active groups. Also similar to the results of earlier studies, orthographic priming produced a switch in the direction of preference between passive and active priming. In contrast, associative priming in the present study produced no such switch; for both the passive and active groups there was a tendency to choose related words. Figure 8 (and the related discussion following Experiment 2, titled ‘ROUSE and Prime Similarity’) explains ROUSE’s ability to predict this complex pattern of

interactions: High values of ρ correspond to orthographic priming and produce predictions as shown in the right hand regions of the upper and lower panels of Figure 8 for passive and active priming. Low values of ρ correspond to associative priming, and produce predictions for passive and active priming as shown in the left hand regions of the upper and lower panels of Figure 8.

ROUSE Predictions for Experiment 6

The parameter estimates for Experiment 6 are given in Table 3. Separate values of ρ are fit to the four conditions of passive and active prime processing crossed with associative and orthographic priming. With only one prime word, the neither-primed and both-primed choice words were necessarily similar to each other. In the model we assumed independent sampling of features resulting in a choice word similarity of ρ^2 for these two conditions.

The predictions given in Figure 13 show ROUSE captures the qualitative trends found in Experiment 6 with the exception of the both-primed benefit with passive associative priming. This benefit is outside the scope of ROUSE and will be taken up in the section titled, '*Short-term Priming Theories and Perceptual Effects*' in the General Discussion.

There are a few quantitative mispredictions in Figure 13. Some of the mismatch might be due to parameter estimates chosen partly to try to capture the both-primed associative benefit (impossible for this version of ROUSE). In addition, as in Experiment 3, we suspected subject differences may have contributed to the quantitative mispredictions, and explored model extensions in which β was sampled from a distribution of values. This

change improved the quantitative predictions, but will not be discussed further in this thesis (for the reasons enunciated in the ROUSE predictions for Experiment 3).

Experiment(s) 7: Compound Word Priming for Constituents

This section consists of five perceptual identification studies with 2-AFC testing that use compound words as primes. On half the trials, the two constituent words of a compound prime word appear as target and foil choice alternatives (i.e. both-primed) while on the other half of trials, constituents of some other compound word appear as choice alternatives (i.e. neither-primed).

One theoretical issue not addressed directly by any of the previous experiments is feature alignment: In order for prime features to activate the same representation as the target flash (thus causing preference), it is critical that the features are aligned properly. The details of such an alignment process are not addressed in this thesis. However one aspect of this issue can be explored by using primes whose parts, rather than wholes, match the subsequent choice words. For instance, for a both-primed test, the prime TAXICABS is followed by the threshold presentation of CABS and then a choice between TAXI and CABS. If such a presentation mode reduces the success of aligning features, then the ROUSE model would predict generally smaller effects, including less both-primed interference.

Using traditional measures of priming, various results have been found with compound words. In long-term repetition priming, Reinitz and Demb (1994) found prior exposure to compound words facilitated stem-completion for both intact and re-paired constituents

whereas with perceptual identification only intact constituents revealed an advantage. Osgood and Hoosain (1974) found prior exposure to noun phrases, single constituents, and re-paired compounds facilitated later perceptual identification on constituents, whereas prior exposure to intact compounds produced no such advantage. In short-term associative priming tested with lexical decision, Sandra (1990) found semantically opaque compounds were not facilitated by priming a constituent (e.g. prime: BREAD target: BUTTERCUP) however semantically transparent compounds were facilitated (e.g. prime: COFFE target: TEASPOON). These results prompted the use of intact and re-paired compound primes and the separation of compounds into different categories according to the extent of phonemic, morphological, and semantic similarity between the compounds and their constituents.

The reported experiments use only both-primed and neither-primed conditions. Two constituent words appear side by side without a separating space; on half the trials this string comprises a valid compound word and on the other half of trials the string consists of randomly re-paired constituents (except in Experiment 7c which only uses intact compound words). In experiment 7a, 7d, and 7e active priming is realized through a lexical decision on intact vs. re-paired constituents. In other words, the task is to respond whether the prime is a valid compound word.

Experiment 7a

Method

There were 41 participants in Experiment 7a.

Four categories of compound words were used. The first two categories were considered *nominal* in that they were nominally compounds based upon orthography but would not typically be considered compounds. *Nominal* compounds were divided into those that contained less phonemic similarity between their constituents and the whole compound, labeled *orthographic* compounds, and those that were more phonemically similar, labeled *orthographic and phonemic*. An example of an *orthographic* compound is COVERAGE breaking into the constituents COVE and RAGE. An example of an *orthographic and phonemic* compound is PLEASING breaking into the constituents PLEA and SING. The remaining two categories of compounds were what would be considered *regular* compounds. These were divided into compounds whose meaning was *opaque* given the constituents (e.g. BEEFCAKE) and compounds whose meaning was *transparent* given the constituents (e.g. TAXICABS).

There were 48 compounds that fell into these 4 different categories of compounds. Breaking these down by word length, 12 compounds were 6 letters long, 24 compounds were 8 letters long, and the remaining 12 compounds were 10 letters long. In all cases both constituents contained the same number of letters. The design fully crossed number of letters with intact/re-paired, both-primed/neither-primed, and whether the target was the left or right constituent. The 2-AFC was presented as a top/down choice and the position of the target was randomly counterbalanced. The left and right constituents were always presented in their corresponding left/right positions in the compound prime even when the prime was a re-paired compound. In the course of the experiment, all constituents appeared on only one trial.

Active priming was accomplished with a lexical decision on the compound prime word. Following one second of fixation point and one second of blank screen, the compound prime word appeared for one second during which participants were instructed to give a lexical decision response. Specifically they were instructed to hit a “YES” button if the letter string comprised a valid compound word and a “NO” button if the letter string consisted of two words randomly stuck together. At the end of the one second of prime presentation, a tone was played for one second if no response was given. If a response was given within the allotted time, but the response was incorrect, a higher tone was played for one second at the time of response in order to provide feedback. No tones were heard if the correct response was given within the one second interval. Regardless of when or if a response was given, the prime remained on screen for one second. The target flash immediately followed the prime, in the same location, followed by a one second pattern mask (minus the target duration), and then the 2-AFC display. Two different buttons were used in the 2-AFC corresponding to the top and bottom words.

72 additional compound words were used in a single practice session of 36 trials. These trials consisted of 18 intact neither-primed and 18 re-paired neither-primed presentations. Target duration was adjusted on the basis of these practice trials in increments of 29ms. Following practice participants received a single block of 128 experimental trials. All other procedures were same as found in Experiment 1. Specific examples of each of the 16 basic conditions (4 types of compounds by intact/re-paired by neither-primed/both-primed) can be found in Table 8. In actuality there were 32 conditions since each of the 16 basic conditions was run with the left constituent as target as well as with the right constituent as target.

Insert Table 8 and Figure 14 about here

Results and Discussion

The active priming lexical decision task proved to be very difficult with average performance only 70% for responses made within the 1 second time limit. This low performance level partly represents an uncontrolled factor: In randomly re-pairing constituents, occasionally valid compounds are created.

Before analyzing perceptual identification performance as a function of priming condition, type of compound, and intact vs. re-paired, an analysis of the left/right position of the target constituent within the compound prime was performed. This analysis is crucial to theories of feature alignment. If feature alignment occurs solely on the basis of left or right justification, then any effects of the other variables should occur only when the target is in the corresponding left or right position. In other words, left/right target position should interact with the other variables. An ANOVA revealed no main effect of left/right position, $F(1,40) = .07$, $p = .79$, and left/right position did not interact with priming condition, $F(1,40) = .17$, $p = .68$, type of compound, $F(3,38) = 1.80$, $p = .16$, or intact/re-paired compound presentation, $F(1,40) = 1.15$, $p = .29$ (all higher order interactions were non-significant as well). Therefore the results shown in Table 8 and Figure 14 are collapsed across left/right position. Considering the effect of priming condition reported next, this suggests feature alignment is either automatic or some mixture of left and right justified.

Accuracy levels are shown in the appropriate column of Table 8. In keeping with previous experiments using two primes, a single compound prime produced a deficit for repetitions of its constituents, $F(1,40) = 8.61$, $p < .01$. This deficit did not interact with type of compound, $F(3,38) = 1.44$, $p = .25$, or with intact/re-paired compound presentation, $F(1,40) = 2.08$, $p = .16$. There was a main effect of type of compound, $F(3,38) = 5.87$, $p < .005$, but not of intact/re-paired, $F(1,40) = .93$, $p = .34$, and these variables did not interact, $F(3,38) = 1.53$, $p = .22$. The results support those of previous experiments: Repetition priming both alternatives results in a performance deficit that is largely independent of manipulations of phonemic and semantic similarity and is therefore assumed to arise from orthographic similarity. That this deficit occurs when primes are presented as a compound word and left/right target position doesn't matter suggest a relatively automatic feature alignment process. The magnitude of the both-primed deficit is less than that seen in previous experiments. This lessening of the prime interference might be due to inaccuracies in the feature alignment process (a mixture of left/right justified) or the specific manner in which the compounds were presented (only for 1 second in a speeded lexical decision task).

The only interaction that approached significance was between priming condition and intact/re-paired compound presentation. As seen in Experiment 7e, this possible interaction is replicated and deserves closer examination. Collapsing across type of compound (and left/right target position), performance as a function of intact/re-paired compound presentation is shown in Figure 14. For intact compound presentation there is a clear both-primed deficit, $t(40) = 3.24$, $p < .0025$, but for re-paired compound presentation there is no difference with priming condition, $t(40) = 1.33$, $p = .19$.

Explanations for this interaction are not immediately obvious. Possibly prime processing is less extensive for re-paired compound presentation since this condition corresponds to a “NO” lexical decision judgment. Essentially the compound might be quickly dismissed in this condition (although not always given the poor performance levels in lexical decision).

A close examination of Table 8 reveals what might be considered an interaction between nominal vs. regular compounds and priming condition. It appears that the both-primed deficit is larger for regular compounds and this was demonstrated with t-tests. However, as seen in Experiment 7d, this trend is seen to reverse and therefore explanations are not considered.

Experiment 7b

Experiment 7a found a both-primed deficit with compound prime presentations under conditions of active priming. Experiment 7b examines whether this deficit exists with passive priming.

Method

There were 51 participants in Experiment 7b. Experiment 7b was identical to Experiment 7a except the lexical decision task was dropped. Compound prime words were still presented for one second. Target duration was determined in increments of 14ms.

Results and Discussion

An ANOVA revealed no main effect of left/right target position, $F(1,50) = .38$, $p = .54$, and no interaction between left/right target position and priming condition, $F(1,50) = .58$, $p = .45$, or intact/re-paired compound presentation, $F(1,50) = .45$, $p = .51$. There was an interaction between left/right target position and type of compound, $F(3,48) = 2.90$, $p < .05$, but this most likely reflects item differences since left and right constituents are separate pools of words. The accuracy results collapsed across left/right target position are shown in the appropriate column of Table 8. There was no main effect of priming condition, $F(1,50) = .48$, $p = .49$, and no interaction between priming condition and type of compound, $F(3,48) = .62$, $p = .61$, or intact/re-paired compound presentation, $F(1,50) = .55$, $p = .46$. Unlike previous experiments, passive priming completely eliminated the both-primed deficit.

If passive priming causes less prime interference, the relatively small both-primed deficit seen in Experiment 7a might have been lessened to the point of being unobservable given the power of this experiment. If this were true, why wasn't a consistent trend of diminished both-primed deficit with passive priming observed in previous experiments? The answer might lie in the specific prime durations used in each experiment. In previous experiments prime duration was not controlled since participants had as much time as needed to respond. As a result, primes were presented for several seconds in active priming whereas the primes were typically presented for only half a second with passive priming. If shorter prime durations produce greater both-primed deficits (up to some limit: near-threshold prime durations would not be expected to produce as much interference) but passive priming diminishes both-primed deficits, these two factors might have offset in

previous experiments. Prime duration is kept constant across Experiments 7a and 7b and this allowed the effect of active vs. passive prime processing to be revealed independent of prime duration.

Some support for idea of increased prime interference with shorter prime durations can be found in Experiment 6. The passive priming condition of Experiment 6 used a 250ms prime duration and the magnitude of the both-primed deficit with orthographic priming is sizable (especially considering the small both-primed deficit seen in Experiment 2 with passive orthographic priming). Additionally, given the assumption in ROUSE of source confusion, there is face validity to the idea of increased interference for shorter prime durations: Shorter durations make it more difficult to visually parse the prime and target and lead to more source confusion.

Experiment 7c

Experiment 7c is a further test of passive priming with compound primes. Unlike Experiment 7b, only intact compound primes are presented. In addition, the specifics of the visual sequence are changed. These changes are in accord with the results of Humphreys, Besner, & Quinlan (1988). They found a deficit in repetition priming as measured with the traditional form of perceptual identification with 300ms of prime duration. The task was to identify the upper case target and primes were always presented in lower case. A secondary task was to write down the identity of the primes (similar to our active priming manipulations). On the assumption that the priming deficit was due to an inability to visually parse the prime and target, they modified prime presentation by presenting primes for 180ms followed by 120ms of pattern mask prior to the target flash.

In addition the secondary task was dropped. These changes caused the priming deficit to change to a priming advantage. Experiment 7c adopts the presentation sequence of Humphreys et al. in order to see if a similar change to a positive repetition effect occurs when preference is controlled.

Method

There were 21 participants in Experiment 7c. The design was similar to Experiment 7b except in the following aspects. Only intact compound prime words were presented doubling the number of data points per condition per subject. Following 250ms of fixation point and 250ms of blank screen, primes were presented for 200ms followed by 100ms of pattern masking and then the target flash. The target was subsequently backward masked by a 250ms pattern mask (minus the duration of the target). Target duration was determined in increments of 8ms as determined by 72 practice trials (all practice trials were both-primed). Compound prime words were presented in lower case, but all other presentations were in upper case.

Results and Discussion

An ANOVA revealed no main effect of left/right target position, $F(1,20) = .84$, $p = .37$, and no interaction between left/right target position and priming condition, $F(1,20) = .64$, $p = .43$, or type of compound, $F(3,18) = .08$, $p = .97$. The accuracy results collapsed across left/right target position are shown in the appropriate column of Table 8. There was no main effect of priming condition, $F(1,20) = .57$, $p = .46$, and no interaction between priming condition and type of compound, $F(3,18) = 1.76$, $p = .19$. Adopting the

presentation sequence that produced a change from negative to positive repetition priming with the traditional form of perceptual identification (Humphreys et al., 1988) did not produce such a change when preference was controlled with both-primed 2-AFC. In light of the discussion associated with Experiment 7b, it might be thought that the shorter prime duration of this experiment would produce a greater both-primed deficit. However such an effect might have been offset by the placement of an intervening mask between the prime and target making it easier to visually separate prime from target.

Experiment 7d

So far all the reported experiments have only looked at short-term priming. In long-term repetition priming, Ratcliff and McKoon (1997) fail to find either a both-primed deficit or advantage. In their experiments several minutes and many intervening words separate study of the primes from their subsequent tests as targets and foils. In our experiments we consistently find a both-primed deficit with short-term repetition priming. It is not clear how much delay between study and test is necessary before the both-primed deficit disappears. If, as assumed in ROUSE, the both-primed deficit is due to source confusion, then the both-primed deficit should disappear with relatively short delays. Experiment 7d tests this theory by repeating Experiment 7a, in which a both-primed deficit was observed, but instead of the compound primes appearing immediately prior to the target flash, the repeated constituents are presented as compound primes on the trial prior to their appearance as target and foil (i.e. lag one priming).

Method

There were 38 participants in Experiment 7d. The design was similar to Experiment 7a except in the following aspects. Following 250ms of fixation point and 250ms of blank screen, upper case compound primes were presented until a non-speeded lexical decision response was given. The lexical decision was non-speeded in order to maximize any potential effect the prime might have on the subsequent trial. As in Experiment 7a, error feedback was given in the form of a tone for incorrect responses. After a lexical decision was made, the compound prime was followed by 250ms of pattern masking and then the target flash. This intervening pattern mask was supposed to provide an opportunity to prepare for the target flash. In this experiment the repeated constituent words were seen in the compound of the previous trial so source confusion is not an issue and an intervening mask should be irrelevant to any priming effects. The target was subsequently backward masked by a 250ms pattern mask (minus the duration of the target). Target duration was determined in increments of 8ms as determined by 72 practice trials (all practice trials were lag one both-primed except the first). All presentations were in upper case.

Results and Discussion

Participants took an average of 950ms (which was therefore the average prime duration) responding to the compound primes and did so with an average accuracy of .834. This is substantially better performance than occurred in Experiment 7a suggesting the primes were highly processed.

An ANOVA revealed a main effect of left/right target position, $F(1,37) = 11.38$, $p < .0025$, as well as an interaction between left/right target position and type of compound,

$F(3,35) = 3.72$, $p < .05$, both of which presumably reflect item differences. Importantly, left/right target position did not interact with priming condition, $F(1,37) = .11$, $p = .74$, or intact/re-paired compound presentation, $F(1,37) = .05$, $p = .83$, and the accuracy results are therefore collapsed across left/right target position as shown in the appropriate column of Table 8.

As in Experiments 7b and 7c, there was no main effect of priming condition, $F(1,37) = .29$, $p = .59$, and no interaction between priming condition and type of compound, $F(3,35) = 2.47$, $p = .08$, or intact/re-paired compound presentation, $F(1,37) = .80$, $p = .38$. As might be expected by the ROUSE theory, the both-primed deficit with active priming seen in Experiment 7a disappears when there is a lag of one trial between presentation of constituents and their subsequent reappearance as target and foil alternatives. This result lends some credence to source confusion as a short-term priming mechanism.

Experiment 7e

Given the null findings of Experiments 7b-7d, the both-primed deficit seen in Experiment 7a might seem questionable. Therefore we decided to replicate Experiment 7a with some slight modifications to the display sequence and with different stimuli. A nearly significant interaction in Experiment 7a suggested that the both-primed deficit was larger for *nominal* compounds as compared to *regular* compounds. However there was no difference between the two types of *nominal* compounds and likewise there was no difference between the two types of *regular* compounds. The next experiment compiles new lists of *nominal* and *regular* compounds without regard to phonemic and semantic similarity.

The results of Experiment 7c and 7d suggest that additional cues (visual or temporal) can be used to eliminate source confusion and therefore eliminate priming. This idea is further tested in a second set of conditions in the next experiment that are identical in every respect except for the inclusion of additional visual cues useful for separating the prime and target.

Method

There were 59 participants in Experiment 7d.

New stimuli were created without reference to the phonemic similarity of the *nominal* compounds to their constituents and without reference to the semantic similarity of the *regular* compounds to their constituents. In total 80 *regular* and 80 *nominal* compounds were 8 letters in length and 44 *regular* and 44 *nominal* compounds were 6 letters in length. As before both constituents contained the same number of letters. Unlike previous experiments, the left/right target position was not fully crossed with the other variables and was instead randomly counterbalanced. All the other variables were fully crossed with one another (top/down position of the target in the 2-AFC was randomly counterbalanced as in the previous designs).

The procedures were similar to Experiment 7a except in the following aspects. Following 500ms of fixation point and 500ms of blank screen, compound primes were presented for 1500ms and lexical decision responses could be given during the first 1000ms of prime presentation. The target flash subsequently occurred followed by a backward pattern mask for 500ms (minus the duration of the target). Feedback through tones was given in the same manner as Experiment 7a. Target duration was determined in

increments of 8ms based upon performance on 72 trials. The threshold determination trials were preceded by 24 practice trials. Both the initial practice trials and the threshold determination trials were run without visual separation cues (which are described below). Following the threshold determination trials, a block of 32 experimental practice trials was given in order that participants adjust to variations in the visual sequence appearing with the introduction of conditions with visual separation cues. All trials throughout the experiment drew upon the same stimuli. There were two blocks of 72 experimental trials.

On half the experimental trials all words were presented in upper case Times Roman font and followed the display sequence listed above. The other half of trials added visual separation cues designed to make it easier to visually parse the prime and target. Unlike the upper case Times Roman target, primes were presented in lower case Cambridge font (old English). Furthermore, primes appeared only for the 1000ms during which a lexical decision response could be given. Then the lower case prime was followed by 500ms of pattern masking prior to the upper case target flash (this kept the SOA of priming equivalent with or without visual separation cues). Conditions with and without these visual cues were run with all levels of the other variables. In total this led to 16 conditions (with/without visual separation cues by regular/nominal compounds by neither/both primed by intact/re-paired compound presentation). Specific examples of these 16 conditions can be found in Table 9.

Insert Table 9 and Figure 15 about here

Results and Discussion

Again, the speeded lexical decision task proved to be very difficult with average performance only 67% for responses made within the 1 second time limit.

An ANOVA on the left/right target position could not be performed since left/right position was randomly counterbalanced and there were many missing cells in crossing left/right position with the basic 16 conditions.

The accuracy results are shown in Table 9. There was no main effect of priming condition, $F(1,58) = .89$, $p = .35$, and priming condition did not interact with type of compound, $F(1,58) = .92$, $p = .34$, intact/re-paired compound presentation, $F(1,58) = .24$, $p = .62$, or the inclusion of visual separation cues, $F(1,58) = .57$, $p = .45$. At first glance it would appear that there were no effects of priming (but see below).

There was a main effect for the inclusion of visual separation cues (overall accuracy was higher), $F(1,58) = 18.72$, $p < .001$, and a nearly significant three-way interaction between visual separation cues, priming condition and intact/re-paired compound presentation, $F(1,58) = 3.77$, $p = .057$. It is this three-way interaction that reveals a both-primed deficit thus replicating Experiment 7a as shown in Figure 15 (collapsed across *regular* and *nominal* compounds). Essentially, as in Experiment 7a, there was an interaction between priming condition and intact/re-paired compound presentation, but the interaction only occurred without visual separation cues. Only with intact compound presentation and no visual separation cues was a both-primed deficit observed, $t(58) = 2.10$, $p < .025$. In other words, both the manipulations of presenting primes as re-paired compounds and the inclusion of visual separation cues eliminated the both-primed deficit. Presumably re-paired presentation eliminated the deficit because compound primes were

processed to a lesser extent and visual separation cues eliminated the deficit because prime and target were more readily visually parsed reducing source confusion. The overall increase in performance with visual separation cues lends support to the second claim: The neither-primed condition presumably suffers some source confusion since orthographic overlap was not controlled; therefore visual separation cues would lead to an increase in performance in this condition (as well as an increase in the both-primed condition).

There was a main effect of type of compound (presumably due to item differences), $F(1,58) = 16.72$, $p < .001$, but type of compound did not interact with any of the other variables and was not involved in any higher order interactions. Nevertheless, an examination of Table 9 shows a trend that is in opposition to the trend observed in Experiment 7a: The both-primed deficit for intact compounds without visual separation cues was stronger for the *regular* compounds than it was for the *nominal* compounds.

General Discussion for Experiment 7

In all the versions of Experiment 7, prime strings that were not words (i.e. re-paired constituent words) produced little in the way of effects. Therefore the discussion to follow is restricted to the case of priming by intact compound words. It is not clear why the re-paired compounds failed to produce any priming effects. Since such re-paired compounds are not valid words, participants might have devoted fewer resources to processing them; if so, even in the active priming experiments, the re-paired compounds may have acted like passive primes. Such an account would be consistent with the finding that passive priming generally failed to produce effects. Why passive processing of compound primes should fail to produce both-primed deficits is less clear, unless problems of alignment generally

reduce effect size. Unlike the passive/active manipulations of Experiments 1-3 and 6, the passive/active manipulation in these experiments used a fixed prime duration; this control of prime duration might also account for the change in the magnitude of the both-primed deficit between passive and active.

Experiment 7a and 7e used active short-term priming, and consistent with, but not as large as, the results for repetition priming, both-primed deficits were found. These deficits were independent of the type of compound suggesting orthography was the main component of priming.

Experiments 7b and 7c demonstrated that the both-primed deficits disappeared with various types of passive priming. Considering that the duration of prime presentation was held constant between the active experiment (7a) and the corresponding passive experiment (7b), this suggests the passive/active manipulations of Experiments 1-3 and 6 were due to more than just different amount of prime exposure.

Experiment 7d demonstrated that the both-primed deficit disappeared under active priming when the compound word was presented in the trial prior to the trial in which the constituents were tested. This result suggests that the both-primed deficit truly is a short-term phenomenon. Experiment 7e, besides replicating the basic finding of a both-primed deficit, demonstrated that the both-primed deficit disappeared when additional visual cues were presented to separate the prime from the target flash. The results of 7d and 7e lend credence to the ideas found in ROUSE: Preference and interference result from prime induced activation (short-term) and an inability to discriminate the source of activation (which could presumably be overcome through the use of additional visual cues).

Although the effects were somewhat unreliable across experiments, the simplest interpretation of the results taken together holds that compound word primes act generally like separate word primes, except that the difficulty of alignment reduces the magnitude of the effects, so that only intact compounds and active processing produce significant effects.

General Discussion

For many researchers and theorists, the present findings may suggest a new way of interpreting short-term priming. With a few notable exceptions (Johnston & Hale, 1984; Hockley & Johnston, 1996; Masson & Borowsky, 1998), previous short-term priming studies leave ambiguous whether the effects of priming are due to perception or preference. The best method we have found to distinguish these involves the use of prime related foil items within a forced-choice task. Using such methods, we obtained strong conclusions, and believe the theoretical implications extend to other short-term priming paradigms such as lexical decision, naming, and the usual form of perceptual identification.

Our studies demonstrate that associative primes can produce perceptual facilitation. However, such effects are found less reliably and are smaller in magnitude than preference effects. Preference effects were found for nearly all experiments and types of primes. Therefore it seems likely that the advantage for related targets found in many prior studies is largely due to what we have termed preference effects. These effects are not to be confused with bias in a signal detection sense, and are not to be cast aside as uninteresting;

they are important, highly constraining for theory, and form the basis for the ROUSE model.

It is particularly important that we were able to reverse the direction of preference through the use of passive versus active processing of primes. Interpretation of prior studies is even more difficult because in many cases it is unclear whether the procedures induced participants to process primes actively or passively (in ROUSE terms, over- or under-estimate α). Under such circumstances it is not surprising that overviews of this literature (e.g. Neely, 1991), reveal a complex and often seemingly contradictory array of results.

The Effects of Passive versus Active Prime Processing

At present we do not have the data to explicate fully the nature of the passive/active manipulation. For example, except for Experiment 7 (which lacked the target-primed and foil-primed conditions necessary for determining the direction of preference), all our experiments confound the passive/active manipulation with prime duration. A recent study in our lab, to be reported elsewhere, varied prime duration under passive instructions; the results showed longer prime durations produced effects like those of active priming, but somewhat smaller in magnitude. Interpretation remains ambiguous, because, for example, longer prime exposure might lead participants to engage in more elaborate (i.e. more active) processing of primes.

We surveyed extant models and did not locate any candidates that appeared to have mechanisms capable of predicting the present passive/active differences. In ROUSE we assume that there are multiple possible sources for an activated feature, and the possible

sources must be taken into account when calculating the evidence values that enter into a Bayesian decision process. Under- or over-estimation of the probability that a feature can be activated by the primes is the key factor that allows ROUSE to predict the differences between active and passive prime processing.

In one interpretation consistent with our model, one system (perhaps a version of working memory) has full access to the primes and their features, presumably due to the primes being presented well above perceptual threshold (and perhaps due to their being attended sufficiently well). Only when this system informs the word recognition system concerning the features common to the primes and choices can preference be removed or reversed within the decision process. This reasoning suggests an alternate interpretation of the passive/active manipulation in terms of prime availability. In the passive condition, primes only appear for half a second and participants are not instructed to pay close attention to the primes. As a result, the identities and the individual features of the primes might not always or all be available at the time of the 2-AFC decision. To the extent that the features of the primes are not known, preference removal (i.e. discounting of evidence) cannot take place and a preference for prime related words will result. Thus a different implementation of ROUSE than the one we used could hold that participants are always excessive in their estimation of α , in both active and passive conditions, but features that are shared between the primes and choices are not discounted when they are not noticed to be present in primes. This idea is further discussed in the subliminal priming section below. In any event we lack the experiments at present to separate these two, nearly equivalent, explanations.

Orthography versus Associations

One puzzle in the short-term priming literature has been an almost universal observation of facilitation for semantic and associative priming whereas orthographic, phonemic, and repetition priming often produce facilitation but sometimes produce deficits (e.g. Meyer & Schvaneveldt, 1974; Humphreys, Besner, & Quinlan, 1988; Lukatela & Turvey, 1996; Dominguez & de Vega, 1997). Our implementation of the ROUSE model suggests an answer to this asymmetry of results. Two factors must both be present for the unrelated choice to be favored: First, evidence consistent with the prime must be discounted to a greater degree than optimal (which occurs in our studies in the active priming conditions); second, the similarity of the prime to the choice must be very high. These conditions can both be satisfied for orthographic and repetition priming, but associative and semantic priming necessarily involve lower levels of similarity, at least for the types of tasks we use in this thesis.

To reiterate the role played in ROUSE by similarity of prime to target, consider again Figure 8 and the attendant discussion. Almost all of the studies in the literature correspond to the target-primed condition of Figure 8. For low similarity primes (i.e. low p), corresponding to associative/semantic priming, the target-primed condition is always above the neither-primed condition (i.e. a positive effect of priming) regardless of the passive/active manipulation. Thus associative/semantic priming will always produce facilitation. For high similarity primes (i.e. high p), corresponding to repetition or orthographic/phonemic priming, the target-primed conditions is above the neither-primed condition with passive priming and below with active priming (i.e. priming can be positive

or negative). Thus orthographic/phonemic and repetition priming can lead to results in either direction (or a lack of priming) depending upon the exact procedures used.

"Subliminal Priming"

Partly in order to minimize the use of explicit strategies in short-term priming (such as ‘guessing’ an answer related to a prime), many priming studies use short duration, masked, sub-threshold primes (e.g. Marcel, 1983; Evett & Humphreys, 1981; Lukatela & Turvery, 1994; Perea & Gotor, 1997). In these experiments, participants are not able to identify the prime on most trials, and may not even be able to determine at an above chance level whether a prime was presented. Extending ROUSE to such ‘subliminal’ paradigms is fairly straightforward. If the system cannot identify the presence of a prime then it is unlikely that any discounting of ON features will occur (and it would certainly be difficult to know which features to discount). Without discounting, any extra evidence from the prime, even if quite weak, is certain to tip the scales in favor of prime related words. This does not imply perceptual facilitation occurs with subliminal priming, but rather that preference can play a role with subliminal priming (indeed it might even play a stronger role since discounting cannot take place). Furthermore, although preferential variability (the factor producing the both-primed deficit) might still occur, the amount of such variability would be very small since at most a very small number of features could be activated by the prime. Thus ROUSE applied to the subliminal priming paradigm would almost certainly predict a ‘beneficial’ result from priming (i.e. a positive preferential shift). Since discounting has been eliminated, this benefit should hold for all types of related primes (i.e. high and low similarity).

It has also been puzzling to the present authors that the magnitude of the beneficial effect of subliminal priming is roughly similar to the magnitude of the beneficial effect of above threshold priming. ROUSE provides a possible account of such results. For long duration explicit primes, there is much activation from the primes, but the decision process substantially discounts the evidence from this activation, reducing the size of the beneficial effect (i.e. the size of the positive preferential shift). For short duration subliminal primes, there is little activation from the primes, but there is no discounting, so the net effect might be similar.

Relevant research in evaluative judgments lends some support to ROUSE's interpretation of subliminal priming. Murphy & Zajonc (1993) used near-threshold and above threshold prime pictures with positive or negative affect; such primes were followed by 'Chinese' letters that had to be evaluated on a positive-negative dimension. Near threshold primes led to congruent evaluation, but above threshold primes led to incongruent evaluation. The authors interpret the results in terms of two different neurological-emotional routes. An interpretation with ROUSE is more parsimonious, holding that the change in the direction of priming is a function of prime availability: With above threshold primes, participants can use knowledge of prime features to discount interfering prime activation resulting in a preference reversal; with near-threshold primes, the prime features are often unknown and no discounting takes place resulting in a preference for targets of similar affect. This same pattern of results has been found with orthographic/phonemic priming as measured with naming (O'Seaghdha & Marin, *in press*). Using high frequency beginning-related prime-target pairs (e.g. prime: STORAGE; target: STORY), 400ms prime exposures resulted in increased latencies whereas brief

(57ms) forward masked primes resulted in decreased latencies. The authors interpret the results in terms of facilitation due to shared orthographic/phonemic features for near-threshold primes and phonological competition for above threshold primes. Again, ROUSE might provide a more parsimonious account.

Short-term Priming Theories and Preferential Effects

For many years, spreading activation theories (Collins & Loftus, 1975; Anderson, 1983) and compound-cue theories (Ratcliff & McKoon, 1988; Doshier & Rosedale, 1989) have been the predominant theoretical accounts of short-term priming in lexical decision. These unidimensional theories map a measure such as summed strength or activation to reaction time; they have not been applied to 2-AFC data. In order to predict the positive preference in our passive priming 2-AFC data, the mechanisms that apply to target processing could be applied to foil processing.

In spreading activation theory, a primed foil could also receive pre-activation and the choice made according to which choice word acquired the greater activation. If only one choice were pre-activated, this choice word would tend to be chosen. If both choices were pre-activated, the result would depend on details of the theory; some sort of stochastic activation would be necessary to produce preferential variability (and a corresponding both-primed deficit). The primary difficulty in accounting for our results with spreading activation is the lack of a decision mechanism for reversing preference with active prime processing. In addition, we have evidence (Experiments 2 and 6), that the extent of representational overlap between prime and target plays a crucial role in determining the

magnitude and even direction of preference. It is not clear how a unidimensional construct like strength of activation could handle such effects.

Compound-cue theory shares these difficulties. The theory posits that the prime plus target are used to probe memory. This could be extended to two-prime studies by using both primes plus the target flash to probe memory. It is not clear how to use the theory to deal with 2-AFC testing. One method seemingly in keeping with the spirit of the model, could use the compound cues to prompt a recall process. The recalled word could then be matched to both choice words. If the primes themselves can be recalled, a preference for a primed word would naturally result. Nonetheless this theory has no decision mechanism capable of producing a preference reversal with active priming, and it is also hard to see how the theory could predict the different preference results for varying degrees of prime/target similarity.

A distributed model of short-term priming developed by Masson (1995; 1991), has been extended to the same/different testing procedure (i.e. is the response word the same or different from the target flash). This model assumes a Hopfield (1982) network of word recognition; every word has a multidimensional energy well learned through Hebbian connections. The effect of a prime is to place the network closer to the location of the target word. Once the target is presented, the system has less distance to travel to settle into the target word attractor. In this manner facilitation is predicted in reaction time data. Masson and Borowsky (1998) extended the model to same/different data by assuming the pattern of activation just prior to the presentation of the response word is stored and then compared to the pattern of activation elicited by the response word. It would be easy to

apply this model to 2-AFC data by comparing the stored pattern to each choice word and choosing the word with the best match.

Similar to ROUSE, Masson's theory assumes a lack of knowledge for the source of activation. Thus the stored pattern will reflect the prime as well as target activation and this would produce a preference for prime related words. Unlike spreading activation and compound-cue theory, this distributed account contains the distributed representation necessary to produce preferential differences with the extent of prime/target overlap. However, in order to produce preference removal or reversal, Masson's theory would require a decisional mechanism similar to that found in ROUSE. If such a mechanism were implemented the resultant theory would be some combination of the two theories: the decisional aspects of ROUSE and the dynamics and learning associated with Masson's model. However, the proper Bayesian calculation for discounting in Masson's model would be very complex; ROUSE was formulated with its ON/OFF activation rule in order to simplify the situation.

Short-term Priming Theories and Perceptual Effects

We have noted that ROUSE cannot predict the (small) both-primed benefit with associative passive priming. We take this result to imply that the prime interacts with and improves target perception. Although current theories do not seem to have mechanisms capable of explaining our preference findings, they tend to include mechanisms producing prime-target interactions, raising the possibility that they could be adapted to provide a mechanism for explaining the benefits of associative priming on target perception.

Most accounts of spreading activation theory are not sufficiently specified to determine whether prime induced target activation is purely additive or interactive. In contrast, compound-cue theory is necessarily interactive; the match between target and each memory trace multiplies the match between prime and each memory trace. If the weighted match values are greater than one, compound-cue theory predicts facilitation above and beyond an additive effect of the prime.

Masson and Borowsky (1998) performed an experiment similar to our Experiment 4 and found, as we did, a both-primed benefit with associative priming (they used a same/different judgment and calculated the A' measure of sensitivity). Masson's distributed model of short-term priming (Masson 1991; 1995) predicted this benefit; the authors define perception in such a way that the gain is not, in their terms, perceptual. In their theory the lower level features that are fed into the net before settling occurs are not affected by the prime. However, we define a perceptual benefit as an interaction between prime and target such that more or less evidence is obtained from the target presentation. In Masson's attractor model it is true that the input is not affected by priming, however the orthographic/phonemic/semantic feature activations interact with target presentation. We would term such an interaction perceptual.

Regardless of the terminology, Masson's distributed model predicts an advantage with priming above and beyond the additive effect of the prime; the prime moves the system to a position near the target, but this advantage does more than provide a fixed advantage to all similarly related words. Once the target is presented, the model settles into the target's attractor more quickly (i.e. more is gained from the target flash). One possible extension of ROUSE capable of predicting a both-primed benefit would be to

take a dynamic approach to prime/target processing such as in Masson's distributed model. This would have the advantage of making explicit predictions about priming as a function of the temporal sequence of events. Nevertheless, as mentioned previously, the proper discounting equations become overly complicated if more sophisticated dynamics are introduced.

Whether in Masson's framework or otherwise, it is important to ask why the perceptual benefit should be specific to associative priming. Alternatively, if a perceptual benefit exists for repetition priming but is outweighed by preferential variability, then it is important to ask why the perceptual benefit should be proportionately larger for associative priming. To explore this question we performed simulations of ROUSE with a simple modification. The target flash activation parameter, β , was multiplied by some ratio greater than one for all features shared by prime and target, in all conditions in which the target was primed. With this addition, and with new estimates of parameter values, ROUSE comes close to predicting the same patterns of results that were predicted by the original version. Unfortunately, the new version is not able to predict simultaneously the both-primed deficit with orthographic priming and the both-primed benefit with associative priming (both of these results were found with the passive group in Experiment 6). If this version of ROUSE is required to predict a both-primed deficit for one level of prime similarity, then it necessarily predicts a both-primed deficit for all levels of prime similarity. On a feature-by-feature basis it amounts to which force is stronger: preferential interference or perceptual gain. If a both-primed deficit is observed, then preferential interference is stronger and ROUSE predicts this will be true for all types of priming (i.e. all values of ρ). One could assume that the perceptual multiplier is higher for associative

features, but this assumption does little more than repeat the data, and does not provide a very satisfying answer to the question.

A more informative explanation involves a dissociation between word level and feature level effects. Suppose perceptual enhancement arises largely because of pre-activation at the word level (perhaps due to top-down support), and not because of activation of lower level features. In fact the lower level activations harm performance due to preferential variability. For repetition priming there is a perceptual gain due to word level effects but this is strongly outweighed by the harm caused by variable activation of all of the orthographic and other low level features. For orthographic priming there is no perceptual gain but the amount of harm is somewhat lower than in the repetition priming case because there is less orthographic feature overlap. The net effect could make the amount of both-primed deficit similar in these two cases. For associative priming, there is no harm caused by variability of orthographic/phonemic feature activation, and there might be enough word level activation to produce the perceptual gains that are seen (i.e. the both-primed benefit). These and related issues would have to be investigated in future research.

Repetition Blindness

The large both-primed deficits with repetition priming are similar in some respects to the phenomenon of repetition blindness (Kanwisher, 1987). In the repetition blindness paradigm, participants view a Rapid Serial Visual Presentation (RSVP) of words and attempt to determine which of the words was repeated. Repetition blindness is the failure to detect the repeat of a word for some period of time following its first presentation. Essentially, participants are retrospectively asking the following question about each word

contained within the RSVP stream: Did this word appear once or twice? This is nominally the same question that applies to our both-primed condition with repetition priming. Most likely our participants realize that both words appeared as primes; the question is which of the two words was repeated as the target. For the repetition blindness paradigm it is an *n*-choice between *n* words that have all been “primed” (i.e. seen at least once recently) and in our experiments it is a two alternative forced choice between two words that have been primed. With the number of primed alternatives greater than two, the repetition blindness paradigm can be seen as maximizing the deficits due to preferential variability.

The ROUSE theory and our present results may shed some light on the competing explanations of repetition blindness (e.g. Whittlesea, Dorken, & Podrouzek, 1995; Downing & Kanwisher, 1995), since the ROUSE theory adopts aspects of both of the predominant accounts. In Kanwisher’s type/token model (1987), repeated words access the same type but there is a failure to individuate repetitions as separate tokens. In terms of tokens there is a perceptual deficit for the second occurrence. ROUSE employs a similar idea as applied at the feature level; similar to tokens there is a failure to individuate the sources of activation for a given feature. The same feature might be activated by both prime and target (i.e. the first and second repetition), but the second activation has no additional effect since activation has a non-linear upper bound. Nevertheless, some other system (corresponding to types) is aware of the identity of the prime and can comprehend the first presentation. The repetition blindness account of Whittlesea, Dorken, and Proudsek (1995) appeals to retrieval interference; both repetitions are contained within memory, but there is a failure to retrieve distinctive information about each. Similarly, in the ROUSE model, the target activates some features while others are activated by the

prime and the decisional mechanism cannot differentiate between these features. This lack of differentiation results in interference, reducing performance. In the future, applying the ROUSE model to repetition blindness data might provide some insight into these and other alternative accounts.

Long-term Repetition Priming

In long-term repetition priming a preference for repeated words has been observed with similar but not dissimilar choice words (Ratcliff & McKoon, 1997). Contrary to this result, Bowers (1999) finds a preference for repeated words regardless of the similarity between choice words. Our Experiment 3 finds each of these patterns, one for active priming, and the other for passive priming. If extended to the long-term priming domain, the decisional mechanism contained within ROUSE could provide an explanation for the conflicting results. One could argue that for some reason (perhaps the nature of the instructions) participants in Ratcliff and McKoon's studies properly estimate the effect of previously studied words (similar to our active priming participants) whereas the participants in Bowers' studies underestimate the effect of previously studied words (similar to our passive priming participants). In an unpublished manuscript, Ratcliff and McKoon (1999) report an experiment testing the effect of instructions on long-term repetition priming. Participants told to passively read a study list displayed a preference for repetitions regardless of choice word similarity whereas participants told to actively study the list of words for a later memory test displayed a preference for repetitions only when the choice words were orthographically similar. Ratcliff and McKoon interpret this difference in terms of a strategy to choose repeated words given the passive study

instructions versus the normal workings of the word recognition system with the active study instructions. Experiment 3 of this thesis provides a short-term version of this instructional difference, and also demonstrates results that shift; perhaps the ROUSE account for the results of Experiment 3 could provide some insight into the differing results found in the long-term priming situation.

The effect of prime presentation in ROUSE is activational and presumably decays with delays between prime presentation and perceptual identification (see Experiment 7). As such, the model does not apply to long-term priming. To produce long-term priming, an additional mechanism is needed to reinstate activation for previously seen features. The mechanism employed by Schooler, Shiffrin, and Raaijmakers (1997) is context matching. The lexical/semantic code for a word is updated with current context features when it is first studied. At test, the choice words have current context features added to their representation, thereby improving the overall match for previously studied words. If the Schooler et. al. theory could be modified to include differential discounting of evidence from context features, depending on the instructions, it might be possible to explain the conflicting results of Ratcliff and McKoon versus Bowers.

Is the ROUSE Theory Too Powerful?

We intentionally chose the simplest possible version of the ROUSE theory, and did not try to augment it with more sophisticated mechanisms and additional processes that would probably make the model more cognitively plausible. The theory can therefore be thought of as a demonstration proof of the power of the core assumptions to predict the findings, and as a stand-in for a class of more complicated models that would incorporate the same

processes. Nonetheless, the success of the ROUSE model, particularly in accounting for data that at first glance appears incomprehensible (and that the authors initially believed disproved the theory), might lead one to question whether the model is too powerful (i.e. too complex). This issue has recently (Myung, Foster, & Browne, 2000) been addressed in the ongoing pursuit of an error measure reflecting model complexity as well as quantitative error. In assessing model complexity, one factor is the number of free parameters (ROUSE is fine on this dimension because it was implemented with few free parameters). A second and critical factor is the proportion of data space that can be predicted by the free range of the parameters. In other words, is the model capable of predicting anything or is it limited to a specific subset of predictions. If a small number of free parameters is nevertheless capable of predicting almost any data pattern, then the model is too complex and therefore untestable. Such complexity measures are still in their infancy and so we have chosen to address this issue in other ways.

In this thesis we have done our best to explain exactly why ROUSE makes the specific predictions that it does. The curious changes in preference as a function of prime and choice word similarity are the natural result of the model under an appropriately sized vector of features and with minimal visual noise. The predictions of the model, although matching the results, were in fact in many instances unanticipated by the authors. Nevertheless the skeptical reader might still worry about model complexity. We address a small part of this concern in the following manner. The parameter estimates reported in Table 3 are the best fitting parameters, tailored as appropriate to each study on the basis of different subjects, different stimuli, and different procedures. However we found that it was possible to use a set of default parameter values⁷ for all experiments and still capture

the correct qualitative pattern of results. This observation suggests that the ROUSE model fits the data for reasons inherent in the basic structure of its assumptions.

Extensions of ROUSE to Other Short-term Priming Tasks

Although ROUSE was developed to deal with performance accuracy in perceptual identification tasks with forced choice testing, it is not hard to envision how its principles could be incorporated in models of other short-term priming tasks. Such extensions are left for future research, but a few brief comments are appropriate in the present setting. Perceptual identification tasks also use same/different testing, and naming. For same/different testing, the pattern of feature activations for the single test item can be used to provide evidence, with evidence discounting proceeding on the basis of feature overlap with the primes much as it is done in the forced-choice setting. For naming all items in the lexicon are potential responses, the evidence for each response being assessed in parallel, the evidence again being modified according to feature overlap with the primes. In order that the system not always emit a response, some minimal level of evidence would be required. A similar approach could be used for other identification tasks such as stem completion, as long as accuracy was the measure of interest.

A number of tasks, most notably lexical decision and naming, involve high levels of accuracy, and response time is the measure of interest. The present approach could be used in such situations if feature activations evolve over time. The evidence evaluation that forms the basis of the ROUSE theory could be carried out continuously. As features activate, a decision would be made at the point in time when sufficient evidence was accumulated to guarantee a high enough level of accuracy.

Conclusions

Through the use of four priming conditions and 2-AFC testing in a perceptual identification task, preferential and perceptual priming effects were distinguished. We verified the existence of perceptual enhancement in the case of associative primes viewed passively, but the effect is quite small and difficult to obtain reliably. On the other hand, preferential effects are large and ubiquitous; the size of such effects increases in magnitude across associative, orthographic, and repetition priming. Preferential effects include at least two components: an increase in variability due to priming that reduces performance, and activation of features due to priming that can produce preferences for or against a choice word related to a prime. When primes were processed passively, the preference was always in favor of prime related words. When primes were processed actively, for a considerable period of time, the preference was much smaller and even reversed for orthographic and repetition priming. These patterns suggest the facilitation typically reported in the literature for short-term priming after passive prime viewing is largely due to a preferential effect rather than a perceptual one. Because our results show that preference can exist in either direction depending upon the manner in which primes are viewed, short-term priming results would be difficult to assess without the use of our control conditions. Without control of preference, subtle differences between paradigms could lead to differences in the magnitude and even direction of priming.

We interpret our results in terms of a theory ‘Responding Optimally with Unknown Sources of Evidence’ (ROUSE). This theory does not yet attempt to incorporate mechanisms for perceptual enhancement due to primes, and therefore cannot and does not

explain the (small) both-primed benefits we found in several studies. In future research such mechanisms will be appended to the theory. It is used instead in this thesis to explain the (large) preference effects found in all the studies. The theory supposes features of the choice words are activated by the primes, by the target, and by visual noise, but the participant is unsure which source(s) activated a given feature. Given this to be the case, the participant discounts the evidence provided by activated features that could have arisen from primes. Depending on the level of discounting, a level that is assumed to vary with active and passive priming, this simple theory explains a wide range of preferential effects and interactions, both positive and negative.

References

- Anderson, J. R. (1983). *The architecture of cognition*. Cambridge, MA: Harvard University Press.
- Becker, S., Moscovitch, M., Behrmann, M., & Joordens, S. (1997). Long-term semantic priming: A computational account and empirical evidence. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, 23, 1059-1082.
- Bowers, J. S. (1999). Visual word priming often reflects perceptual learning rather than perceptual bias: Reply to Ratcliff and McKoon (1997). *Manuscript submitted for publication*
- Bowers, J. S., & Michita, Y. (1998). An investigation into the structure and acquisition of orthographic knowledge: Evidence from cross-script Kanjii-Hiragana priming. *Psychonomic Bulletin & Review*, 5, 259-264.
- Collins, A. M., & Loftus, E. F. (1975). A spreading-activation account theory of semantic processing. *Psychological Review*, 82, 407-428.
- Coltheart, M. (1980). Iconic memory and visible persistence. *Perception and Psychophysics*, 27, 183-228.
- Curran, P. J., West, S. G., & Finch, J. H. (1996). The robustness of test statistics to nonnormality and specification error in confirmatory factor analysis. *Psychological Methods*, 1, 16-29.
- Dominguez, A., & de Vega, M. (1997). lexical inhibition from syllabic units in spanish visual word recognition. *Language and cognitive processes*, 12, 401-422.
- Dosher, B. A., & Rosedale, G. (1989). Integrated retrieval cues as a mechanism for priming in retrieval from memory. *Journal of Experimental Psychology: General*, 118, 191-211.

- Downing, P. & Kanwisher, N. (1995). A reply to Whittlesea, Dorken, and Podrouzek (1995) and Whittlesea and Podrouzek (1995). *Journal of Experimental Psychology: General*, 121, 1698-1702.
- Evett, L. J. and Humphreys, G. W. (1981). The use of abstract graphemic information in lexical access. *Quarterly Journal of Experimental Psychology*, 33, 325-350.
- Hochhaus, L., & Johnston, J. C. (1996). Perceptual repetition blindness effects. *Journal of Experimental Psychology: Human Perception and Performance*, 22, 355-360.
- Humphreys, G. W., Besner, D., & Quinlan, P. T. (1988). Even perception and the word repetition effect. *Journal of Experimental Psychology: General*, 117, 51-67.
- Hogben, J. H., & Di Lollo, V. (1974). Perceptual intergration and perceptual segregation of brief visual stimuli. *Vision Research*, 14, 1059-1069.
- Hopfield, J. J. (1982). Neural networks and physical systems with emergent collective computational abilities. *Proceedings of the National Academy of Sciences U. S. A.*, 79, 2554-2558.
- Johnson, M. K., Hashtroudi, S., & Lindsay, D. S. (1993). Source monitoring. *Psychological Bulletin*, 114, 3-28.
- Johnston, J. C., & Hale, B. L. (1984). The influence of prior context on word identification: Bias and sensitivity effects. In H. Bouma & D. G. Bouwhuis (Eds.), *Attention and performance X: Control of language processes* (pp 243-255). Hillsdale, NJ: Erlbaum.
- Kanwisher, N. G. (1987). Repetition blindness: Type recognition without token individuation. *Cognition*, 27, 117-143.
- Kucera, H., & Francis, W. N. (1967). *Computational analysis of present-day American English*. Providence, RI: Brown University Press.
- Lombardi, W. J., Higgins, E. T., & Bargh, J. A. (1987). The role of consciousness in priming effects on categorization: Assimilation versus contrast as a function of

- awareness of the priming task. *Personality and Social Psychology Bulletin*, 13, 411-429.
- Lukatela, G., & Turvey, M. T. (1994). Visual lexical access is initially phonological: I. Evidence from associative priming by words, homophones, and pseudo homophones. *Journal of Experimental Psychology: General*, 123, 107-128.
- Lupker, S. J., & Colombo, L. (1994). Inhibitory effects in form priming: evaluating a phonological competition explanation. *Journal of Experimental Psychology: Human Perception and Performance*, 20, 437-451.
- 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.
- Martin, L. L. (1986). Set/reset: Use and disuse of concepts in impression formation. *Journal of Personality and Social Psychology*, 51, 493-504.
- Masson, M. E. J. (1991). A distributed memory model of context effects in word identification. In D. Besner & G. W. Humphreys (Eds.), *Basic processes in reading: visual word recognition*. Hillsdale, NJ: Erlbaum.
- Masson, M. E. J. (1995). A distributed memory model of semantic priming. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 21, No. 1, 3-23.
- Masson, M. E. J., & Borowsky, R. (1998). More than meets the eye: Context effects in word identification. *Memory & Cognition*, 26, 1245-1269.
- McKoon, G., & Ratcliff, R. (1992). Spreading activation versus compound cue accounts of priming: mediated priming revisited. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, 18, 1155-1172.
- McNamara, T. P. (1992). Theories of priming: I. Associative distance and lag. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, 18, 1173-1190.

- McNamara, T. P. (1994). Theories of priming: II. Types of primes. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, 20, 507-520.
- McRae, K., & Boisvert, S. (1998). Automatic semantic similarity priming. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, 24, 558-572.
- Meyer, D., & Schvaneveldt, R. W. (1971). Facilitation in recognizing parts of words: Evidence of a dependence between retrieval operations. *Journal of Experimental Psychology*, 90, 227-234.
- Meyer, D., & Schvaneveldt, R. W., & Ruddy, M. G. (1974). Functions of graphemic and phonemic codes in visual word-recognition. *Memory & Cognition*, 2, 309-321.
- Murphy, S. T., & Zajonc, R. B. (1993). Affect, cognition, and awareness: affective priming with optimal and suboptimal stimulus exposures. *Journal of Personality & Social Psychology*, 64, 723-739.
- Mussweiler, T., & Neumann, R. (in press). Sources of mental contamination: comparing the effects of self-generated versus externally-provided primes. *Journal of Experimental Social Psychology*.
- Myung, I. J., Foster, R. M., & Browne, W. M., eds. (2000). Special Issue on Model Selection. *Journal of Mathematical Psychology*, 44, 000-000.
- Neely, J. H. (1991). Semantic priming effects in visual word recognition: A selective review of current findings and theories. In D. Besner & G. W. Humphreys (Eds.), *Basic processes in reading: Visual word recognition* (pp. 264-336). Hillsdale, NJ: Erlbaum
- Nelson, D. L., McEvoy, C. L., & Schreiber, T. A. (1994). *The University of South Florida word association, rhyme and word fragment norms*. Unpublished manuscript.
- Norris, D. (1995). Signal detection theory and modularity: On being sensitive to the power of bias models of semantic priming. *Journal of Experimental Psychology: Human Perception and Performance*, 21, 935-939.

- O'Seaghdha, P. G., & Marin, J. W. (1997). Mediated semantic-phonological priming: calling distant relatives. *Journal of Memory and Language*, 36, 226-252.
- O'Seaghdha, P. G., & Marin, J. W. (in press). Phonological competition and cooperation in form-related priming: sequential and nonsequential processes in word production. *Journal of Experimental Psychology: Human Perception and Performance*.
- Osgood, C. E., & Hoosain, R. (1974). Salience of the word as a unit in the perception of language. *Perception & Psychophysics*, 15, 168-192.
- Perea, M., & Gotor, A. (1997). Associative and semantic priming effects occur at very short stimulus-onset asynchronies in lexical decision and naming. *Cognition*, 62, 223-240.
- 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. (1988). A retrieval theory of priming memory. *Psychological Review*, 95, 385-408.
- Ratcliff R., & McKoon, G. (1997). A counter model for implicit priming in perceptual word identification. *Psychological Review*, 104, 319-343.
- Ratcliff R., & McKoon, G. (1999). Bias in the counter model. Unpublished manuscript.
- Reinitz, M. T. & Demb, J. B. (1994). Implicit and explicit memory for compound words. *Memory & Cognition*, 22, 687-694.
- Sandra, D. (1990). On the representation and processing of compound words: automatic access to constituent morphemes does not occur. *The Quarterly Journal of Experimental Psychology*. 42A, 529-567.
- Schooler, L. J., Shiffrin, R. M., & Raaijmakers, J. G. W. (1997). A model for implicit effects in perceptual identification. *Manuscript submitted for publication*.
- Schwarz, N., & Bless, H. (1992). Constructing reality and its alternatives: An inclusion/exclusion model of assimilation and contrast effects in social judgment. In L.

- L. Martin and A. Tessel (Eds.), *The construction of social judgments* (pp. 217-245). Hillsdale, NJ: Erlbaum.
- Sperling, G., & Weichselgartner, E. (1995). Episodic theory of the dynamics of spatial attention. *Psychological Review*, 102, 503-532.
- Taylor, G. A., & Chabot, R. J. (1978). Differential backward masking of words and letters by masks of varying orthographic structure. *Memory & Cognition*, 6, 629-635.
- Thompson-Schill, S. L., Kurtz, K. J., & Gabrieli, J. D. E. (1998). Effects of semantic and associative relatedness on automatic priming. *Journal Memory and Language*, 8, 440-458.
- Wegener, D. T., & Petty, R. E. (1995). Flexible correction processes in social judgment: the role of naive theories in corrections for perceived bias. *Journal of Personality & Social Psychology*, 68, 36-51.
- Whittlesea, B. W. A., Dorken, M. D., & Podrouzek, K. W. (1995). Repeated events in rapid lists: Part I. Encoding and representation. *Journal of Experimental Psychology: General*, 121, 1670-1688.

Author Note

This research was supported by PHS-NIMH Grant 12717; NSF Grant SBR 9512089; NIMH Training Grant MH19879. David Huber was supported by an Indiana University College of Arts and Sciences (COAS) Dissertation Year Research Fellowship and a National Science Foundation (NSF) Graduate Research Fellowship (DGE-9253867 006).

Correspondence concerning this thesis should be addressed to David E. Huber, who is now at the Department of Psychology, University of Colorado, Boulder, Campus Box 345, Boulder, Colorado 80309-0345. E-mail may be sent to dhuber@psych.colorado.edu.

Footnotes

1. To a first approximation, one can alternatively characterize the activation of choice features in terms of matching a perceived vector of features to the choice vectors: ON features in a choice would be mapped to features that MATCH the perceived vector, and OFF features in a choice would be mapped to features that MISMATCH the perceived vector.
2. If our model were more thoroughly developed, we would lay out more precisely the factors that ought to affect the values of these three parameters. For example, β might vary with aspects of target presentation and duration, γ might vary with mask and presentation characteristics, and α might vary with time between prime and target flash. Such details are not needed at the present stage of development of the ROUSE model.
3. The concept of discounting may seem clear enough that the reader sees no need for additional explanation. However, the model's predictions for some of the subsequent studies are far from intuitively clear, and the next two sections will help the reader to follow the subsequent developments.
4. In actuality the re-pairing was more complicated than switching the primes between two sets of intact pairs. With a simple switching, subjects could potentially use their memory of previous rearranged trials to discern the answer on later rearranged trials. Instead, a one-offset rearrangement was used (e.g. $B' \rightarrow A$, $C' \rightarrow B$, $A' \rightarrow C$).
5. Subtle differences in instructions can produce alternatives to the Ratcliff and McKoon (1997) pattern of results: Bowers (1999) found strong preferences for prime related choices even when the choices were dissimilar, and also found evidence for enhanced perception due to priming.
6. The mathematically sophisticated reader should find it easy to verify higher performance predictions for target-primed than foil-primed, for the parameters associated with Figure 8, for an extreme case of high choice word similarity such

that only one feature is relevant. It follows that as choice word similarity increases up to this extreme case, there must be a crossover point at which the direction of preference changes. Determining this crossover point requires simulation.

7. These default parameters were the same as those listed with and used in the creation of Figures 6, 8, and 9. The only parameter that cannot be given a default value is prime similarity, ρ . Sensibly, this parameter is highly dependent upon the type of primes used.

Table 1
Experiment 1: Examples and Results

Type of priming and Priming condition	Primes	Target	Choice words	Passive p(c)	Active p(c)
neither	CHEF + ACRE			.692	.760
Associative					
both	SOCK + TOAD			.671	.793
target	SOCK + ACRE			.733	.790
foil	CHEF + TOAD			.670	.776
Repetition			SHOE		
both	SHOE + FROG	SHOE	or	.626	.647
target	SHOE + ACRE		FROG	.770	.716
foil	CHEF + FROG			.567	.781
Repetition (alt. assoc.)					
target	SHOE + TOAD			.712	.686
foil	SOCK + FROG			.594	.720

Note. p(c) = forced-choice performance; alt. assoc. = the alternative choice word was associatively primed.

Table 2
ROUSE Log-Odds Summary Statistics

α'	Priming Condition			
	neither-primed	both-primed	target-primed	foil-primed
0				
M	1.26	1.14	3.69	-1.27
Mdn	1.29	1.29	3.87	-1.29
SD	1.67	2.81	2.27	2.36
.05				
M	1.26	0.50	1.33	0.43
Mdn	1.29	0.57	1.13	0.16
SD	1.67	1.24	1.24	1.66
.3				
M	1.26	0.14	-0.02	1.41
Mdn	1.29	0.16	0.31	1.13
SD	1.67	0.34	0.86	1.47

Note. Statistics are for Figure 6 Distributions.

Table 3
Best fitting ROUSE parameters

Probability parameter	Experiment 1		Experiment 2		Experiment 3		Experiment 4	
	Passive	Active	Passive	Active	Passive	Active	Passive	Active
α (prime actual)	.073	.085	.105	.110	.112	.090	.379	.167
α' (prime estimate)	.054	.152	.075	.152	.097	.125	.290	.999 ^a
β (target flash)	.034	.054	.053 / .077 ^b	.055 / .056 ^b	.046 / .074 ^b	.062 / .083 ^b	.048 / .037 ^b	.062 / .056 ^b
ρ (associative)	.296				---		.078 / 0.0024 ^c	
ρ (orthographic)	---				.700		.854 / .027 ^c	
$\Sigma\chi^2$ (error)	11.05				96.59		12.30	

Note. Dashes indicate inapplicable parameters

^a Simulations revealed that setting the estimate of alpha, α' , to values approaching the actual α level while holding the other parameters constant produced little change in the fit to the observed data. Nevertheless, the parameter estimation routine was able to find miniscule improvements in the fit by increasing α' to its maximum value.

^b Experiments containing separate pools of target words for different conditions were allowed separate target encoding parameters, β s, for each pool of target words. In Experiment 2, β s refer to the orthographic and repetition priming conditions (on the left and right of the forward slash, correspondingly). In Experiment 3, the β s refer to the dissimilar and similar choice word conditions. In Experiment 4, the β s refer to the associative and orthographic priming conditions.

^c In Experiment 4, the neither-primed and both-primed conditions necessarily introduced some degree of similarity between the choice words. Therefore in Experiment 4, the similarity parameter, ρ , to the left of the forward slash, refers to prime similarity, and the second ρ after the forward slash is choice-word similarity.

Table 4
Experiment 2: Examples and Results

Priming condition	Primes	Target	Choice words	Passive p(c)	Pctive p(c)
Type of priming					
Passive: <i>orthographic</i> and <i>orthographic and phonemic</i> ; case switch					
Active: <i>orthographic</i> (examples shown)					
neither	DATA + FLAG			.757	.801
both	AIRY + HALO	AWRY	AWRY	.740	.716
target	AIRY + FLAG		or HALT	.845	.762
foil	DATA + HALO			.628	.763
Type of priming					
Passive: <i>orthographic</i> and <i>orthographic and phonemic</i> ; same case					
Active: <i>orthographic and phonemic</i> (examples shown)					
neither	PIER + COLT			.750	.798
both	HAIL + DUAL	HALE	HALE	.765	.673
target	HAIL + COLT		or DUEL	.824	.750
foil	PIER + DUAL			.668	.743
Type of priming: <i>orthographic and morphologic</i>					
neither	LIFT + DIVE			.791	.807
both	BEND + KNEW	BENT	BENT	.731	.717
target	BEND + DIVE		or KNOW	.867	.766
foil	LIFT + KNEW			.648	.766
Type of priming: <i>repetition</i>					
neither	BELL + KNEE			.833	.780
both	GRIP + JURY	GRIP	GRIP	.699	.655
target	GRIP + KNEE		or JURY	.860	.722
foil	BELL + JURY			.782	.765

Note. p(c) = forced-choice performance

Table 5.

Experiment 3: Examples and Results

choice word similarity	priming condition	primes	target	2-AFC	passive p(c)	active p(c)
passive orthographic active orth/phon case switch	neither	DATA + FLAG	AWRY	AWRY or AIRY	.724	.740
	both	AWRY + AIRY			.609	.613
	target	AWRY + FLAG			.709	.735
	foil	DATA + AIRY			.558	.625
passive +phonemic active orth/phon same case	neither	PIER + COLT	HALE	HALE or HAIL	.694	.740
	both	HALE + HAIL			.564	.607
	target	HALE + COLT			.692	.701
	foil	PIER + HAIL			.529	.537
orthographic +semantic	neither	LIFT + DIVE	BENT	BENT or BEND	.665	.723
	both	BENT + BEND			.633	.606
	target	BENT + DIVE			.750	.726
	foil	LIFT + BEND			.724	.542
neutral	neither	BELL + KNEE	GRIP	GRIP or JURY	.609	.793
	both	GRIP + JURY			.709	.719
	target	GRIP + KNEE			.558	.766
	foil	BELL + JURY			.694	.768

Table 6.

Experiment 5: Examples and Results

priming condition	choice word similarity	prime for neither	prime for primed	target	2-AFC	neither p(c)	primed p(c)
both	dissimilar	MILK	THEME	SONG	SONG or IDEA	.741	.756
	semantic	FLOAT	HUM		SONG or TUNE	.758	.772
	orth+sem				SONG or SING	.718	.743
target	dissimilar	MILK	THEME	SONG	SONG or HOUR	---	.795
	dissimilar	FLOAT	HUM		SONG or LONG	.735	.803
	orthographic				SONG or LONG	.779	.845
foil	dissimilar	FLOAT	HUM	HOUR	HOUR or SONG	.733	.696
	orthographic			LONG	LONG or SONG	.698	.619

Table 7.

Experiment 6: Examples and Results

type of priming	priming condition	prime	target	2-AFC	passive p(c)	active p(c)
forward associative	neither	WOUND	CABIN	CABIN or HOTEL	.749	.797
	both	LODGE			.756	.802
	target	LODGE		CABIN or BEACH	.797	.829
	foil	SHELL			.705	.781
backward associative	neither	BOARD	PSALM	PSALM or VERSE	.735	.811
	both	BIBLE			.755	.810
	target	BIBLE		PSALM or BRAWL	.794	.798
	foil	FIGHT			.692	.771
symmetric associative	neither	WHOLE	HAPPY	HAPPY or FROWN	.730	.824
	both	SMILE			.787	.802
	target	SMILE		HAPPY or ABOVE	.831	.836
	foil	BELOW			.691	.771
orthographic	neither	HEDGE	DUSTY	DUSTY or MISTY	.717	.757
	both	MUSTY			.591	.700
	target	MUSTY		DUSTY or SHAPE	.739	.710
	foil	SHAKE			.618	.779

Table 8.

Experiment 7a-d: Examples and Results

type of compound	presentation	priming condition	prime	target	2-AFC	7a p(c)	7b p(c)	7c p(c)	7d p(c)
orthographic (nominal)	intact	neither	COVERAGE	PEAS	ANTS or PEAS	.732	.676	.815	.763
		both	PEASANTS			.640	.630	.765	.717
	re-paired	neither	COVEANTS		RAGE or PEAS	.716	.650	---	.766
		both	PEASRAGE			.704	.637	---	.734
+phonemic (nominal)	intact	neither	MISSPOKE	SING	PLEA or SING	.744	.689	.783	.717
		both	PLEASING			.655	.676	.786	.783
	re-paired	neither	PLEAPOKE		MISS or SING	.701	.667	---	.714
		both	MISSSING			.643	.662	---	.753
opaque (regular)	intact	neither	BEEFCAKE	NOSE	NOSE or DIVE	.777	.706	.815	.773
		both	NOSEDIVE			.747	.740	.845	.796
	re-paired	neither	BEEFDIVE		NOSE or CAKE	.726	.701	---	.783
		both	NOSECAKE			.741	.674	---	.750
transparent (regular)	intact	neither	TAXICABS	BOAT	BOAT or SAIL	.753	.674	.842	.793
		both	SAILBOAT			.735	.701	.830	.770
	re-paired	neither	SAILCABS		BOAT or TAXI	.744	.694	---	.763
		both	TAXIBOAT			.701	.679	---	.724

Table 9.
Experiment 7e: Examples and Results

type of compound	presentation	priming condition	prime	target	2-AFC	7e p(c)
nominal (replication)	intact	neither	COVERAGE	PEAS	ANTS or PEAS	.685
		both	PEASANTS			.657
	re-paired	neither	COVEANTS		RAGE or PEAS	.665
		both	PEASRAGE			.659
nominal (visual separation)	intact	neither	coverage	PEAS	ANTS or PEAS	.731
		both	peasants			.734
	re-paired	neither	coveants		RAGE or PEAS	.750
		both	peasrage			.699
regular (replication)	intact	neither	BEEFCAKE	NOSE	NOSE or DIVE	.751
		both	NOSEDIVE			.693
	re-paired	neither	BEEFDIVE		NOSE or CAKE	.693
		both	NOSECAKE			.710
regular (visual separation)	intact	neither	beefcake	NOSE	NOSE or DIVE	.746
		both	nosedive			.763
	re-paired	neither	beefdive		NOSE or CAKE	.753
		both	nosecake			.766

Figure Captions

Figure 1. The sequence of displays for trials in most experiments. The display contained within the dashed box only appears in the active priming condition. Prior to this sequence, a fixation point followed by a blank screen are each displayed for 500ms. Presentations are synched to the vertical refresh of a PC monitor running at 120Hz. An initial block of trials progressively determines an appropriate target flash duration such that participants are 75% correct.

Figure 2. The theory, Responding Optimally with Unknown Sources of Evidence (ROUSE) assumes three sources independently activate target and foil features. The values, β (target flash), γ (noise), and α (primes), are the probabilities that each source has activated a feature and that the feature remains active until the decision process is initiated. Noise is assumed to arise from the mistaken perception of letters due to the pattern mask. Similarity between prime and choice word (depending upon the condition) is assumed to be zero for unrelated primes, the parameter ρ for related primes, and one for identical primes (i.e. repetition priming). In the simulations, lexical entries are represented by 20 features. These features might contain orthographic, phonemic, or semantic information.

Figure 3. The accuracy results and predictions for Experiment 1. Error bars are two standard errors of the mean. Passive versus active priming is a between subjects manipulation. For comparison to the three types of priming, the single neither-primed condition is displayed three times.

Figure 4. The feature likelihood ratios that might appear in the numerator or denominator product terms of Equation 2. This is a 2 X 2 contingency depending on whether a feature is active or inactive and in the primes or not in the primes. It is assumed that the primes (and their features) are known. The numerator of each ratio is conditional on the feature

existing within the target (and not the foil) and the denominator is conditional on the feature existing within the foil (and not the target). Features that appear in both the target and foil provide no discriminating information and are not considered in the decision process.

Figure 5. Venn diagrams comparing target and foil evidence without discounting and with optimal discounting. The target, noise, and prime(s) regions include features that are active as a result of the labeled source. As seen in Figure 4, only three levels of evidence exist for each feature. OFF features (colored black) provide evidence against the target (ratio less than one) while ON features that are not contained in the prime(s) (colored white) provide strong evidence in favor of the target (ratio greater than one). Discounted features (colored grey) are active features that appeared in the prime(s) and therefore provide weaker evidence in favor of the target. The activation probabilities are low in the simulations, hence the small areas of overlap between target, prime(s), and noise. Additionally the black, OFF feature regions, are in reality much larger. These diagrams show a situation of similarity priming (e.g. associative priming: $0 < \rho < 1$); the proportion of target and noise activated features that appear with a discounted level of evidence (i.e. the grey circles within the white regions) corresponds to the similarity probability, ρ .

Figure 6. The distribution of simulated log-odds (log of Equation 1) in the four priming conditions with repetition priming ($\rho = 1$) for three different levels of discounting ($\alpha = .1$; $\alpha' = 0$, $\alpha' = .05$, and $\alpha' = .3$). The neither-primed condition is only shown once since it is unaffected by discounting. The other parameters are: $N = 20$, $\gamma = .02$, $\beta = .05$.

Figure 7. The accuracy results and predictions for Experiment 2. Error bars are two standard errors of the mean. Passive versus active priming is a between subjects manipulation.

Figure 8. Simulated results as a function of prime similarity (ρ). There is a crossover from a positive preferential shift to a negative preferential shift for the active condition as prime

similarity increases. For the passive condition the preferential shift is always positive. The neither-primed and both-primed conditions are unaffected by discounting and identical in the two panels. Probability correct is averaged across 20,000 simulations for each value of ρ . The default parameters used in the simulations were: $N = 20$, $\gamma = .02$, $\beta = .05$, $\alpha = .1$, α' (passive) = .05, α' (active) = .3.

Figure 9. The distribution of simulated log-odds (log of Equation 1) for the target-primed and foil-primed conditions for three different levels of prime similarity (ρ) with active priming ($\alpha = .1$; $\alpha' = .3$). The other parameters are: $N = 20$, $\gamma = .02$, $\beta = .05$.

Figure 10. The accuracy results and predictions for Experiment 3. Error bars are two standard errors of the mean. Passive versus active priming is a between subjects manipulation.

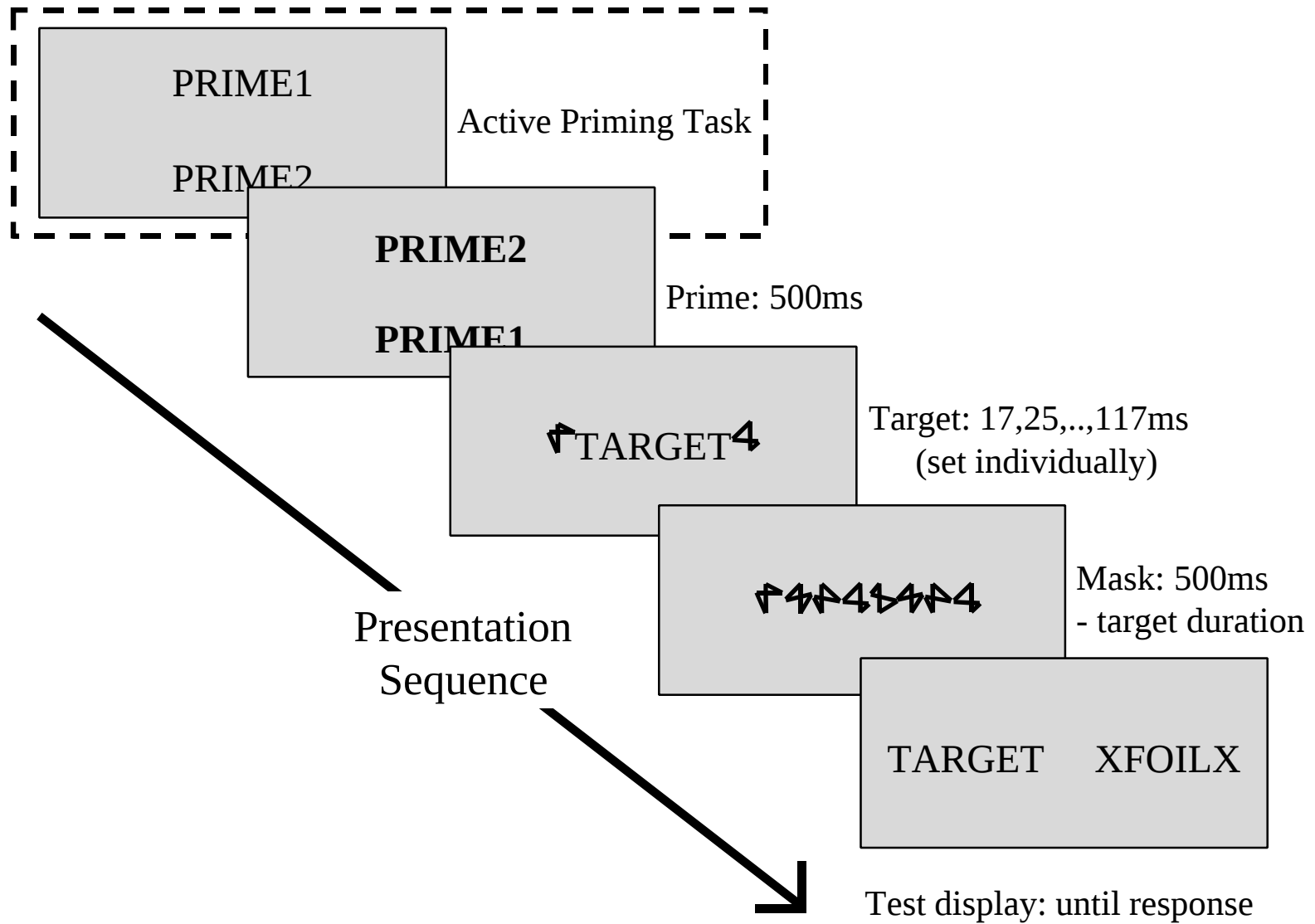
Figure 11. The neither-primed and both-primed accuracy results for Experiment 4. Error bars are two standard errors of the mean.

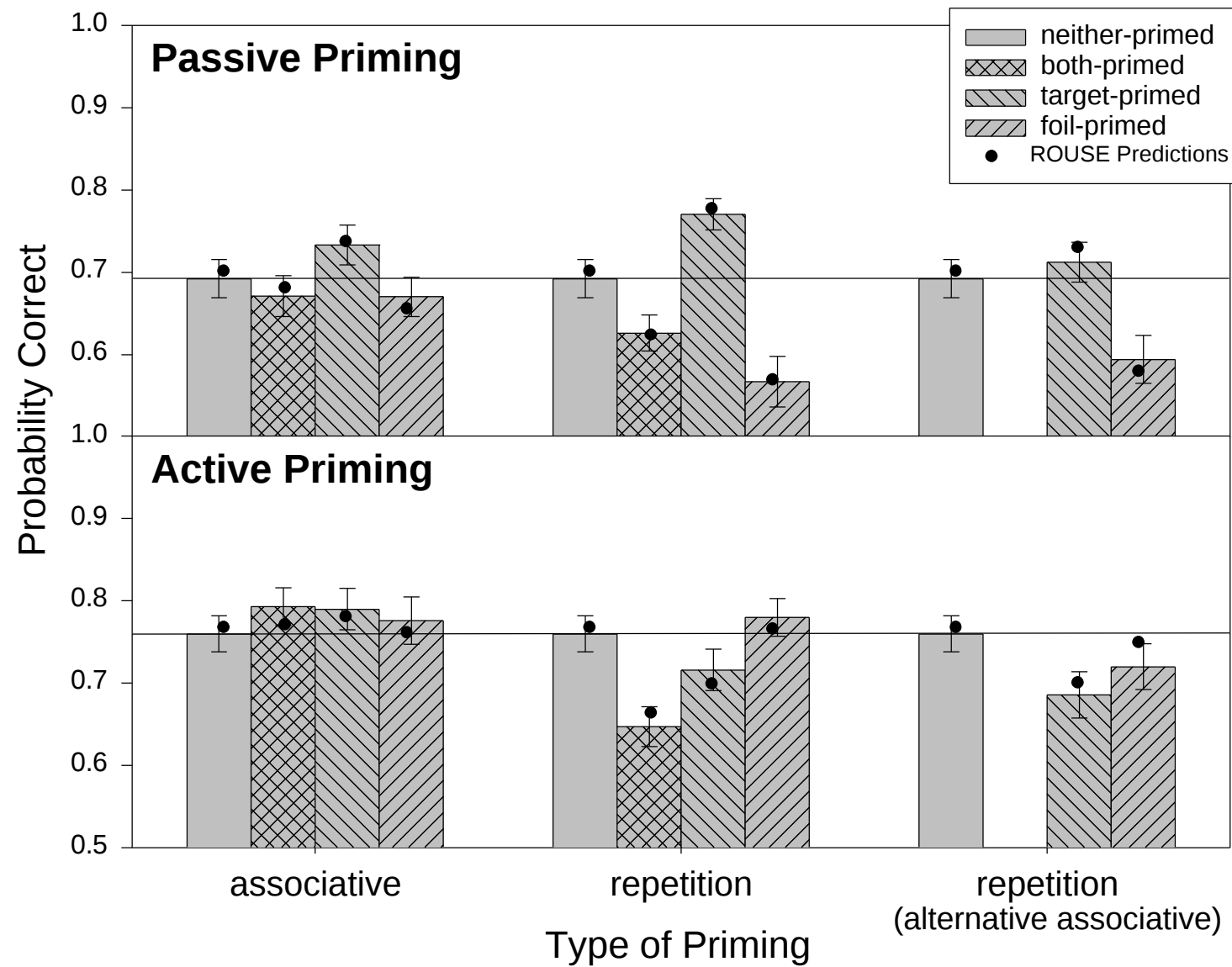
Figure 12. The neither-primed and pattern-mask accuracy results for Experiment 4. Error bars are two standard errors of the mean.

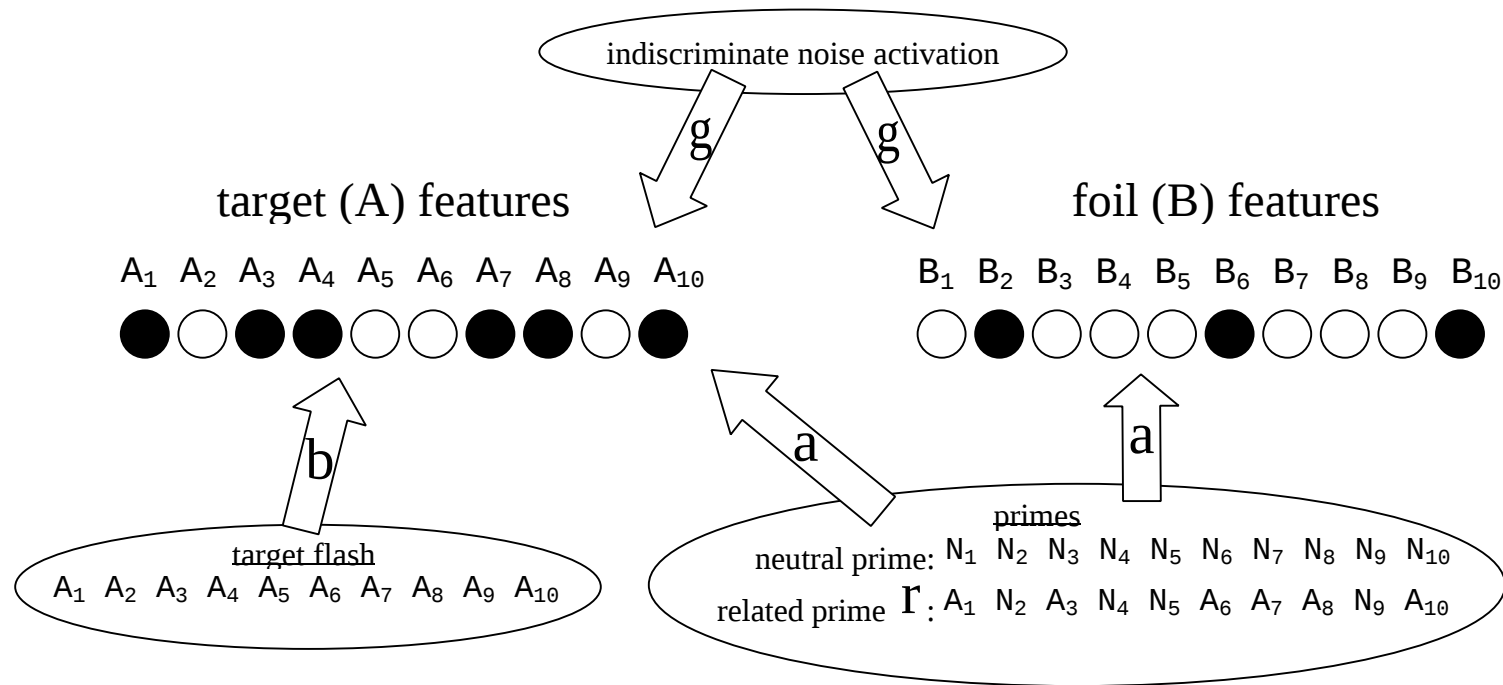
Figure 13. The accuracy results and predictions for Experiment 6. Error bars are two standard errors of the mean. Passive versus active priming is a between subjects manipulation.

Figure 14. The accuracy results for Experiment 7a. Error bars are two standard errors of the mean.

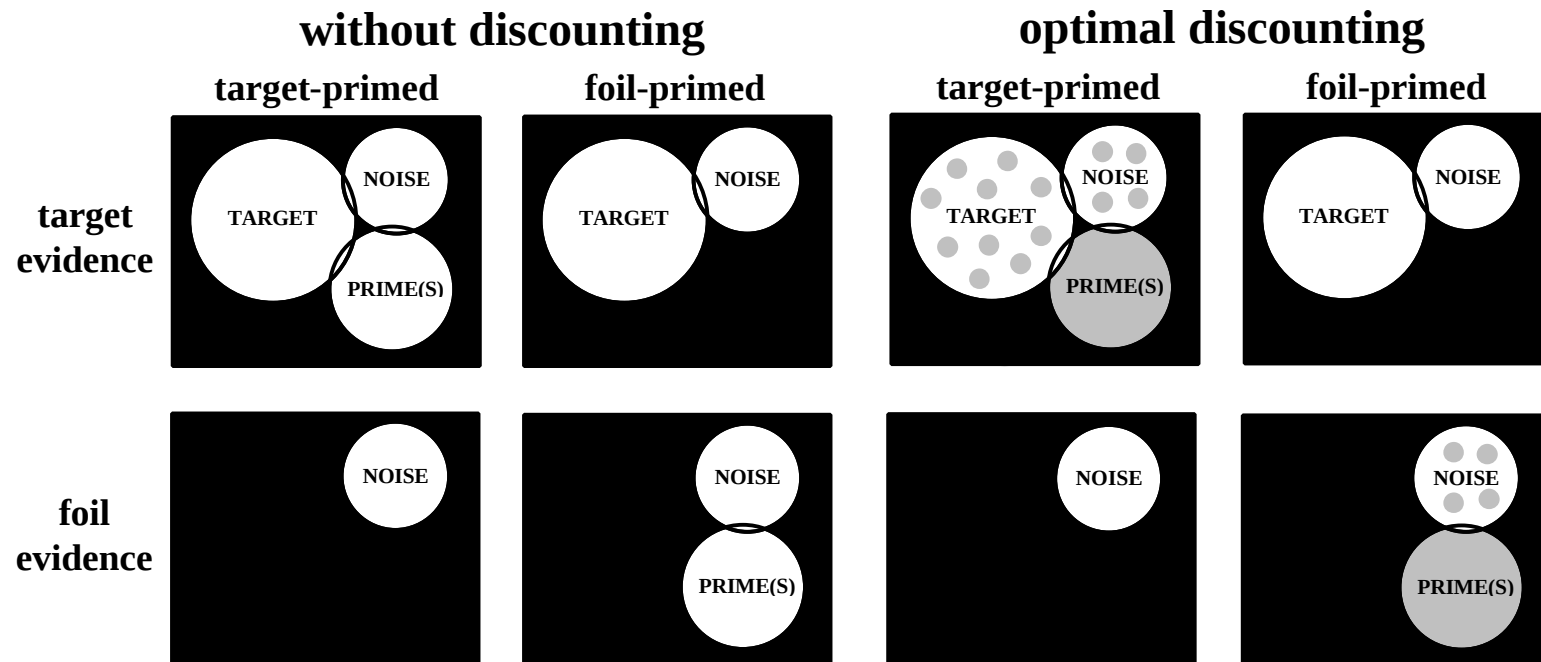
Figure 15. The accuracy results for Experiment 7e. Error bars are two standard errors of the mean.

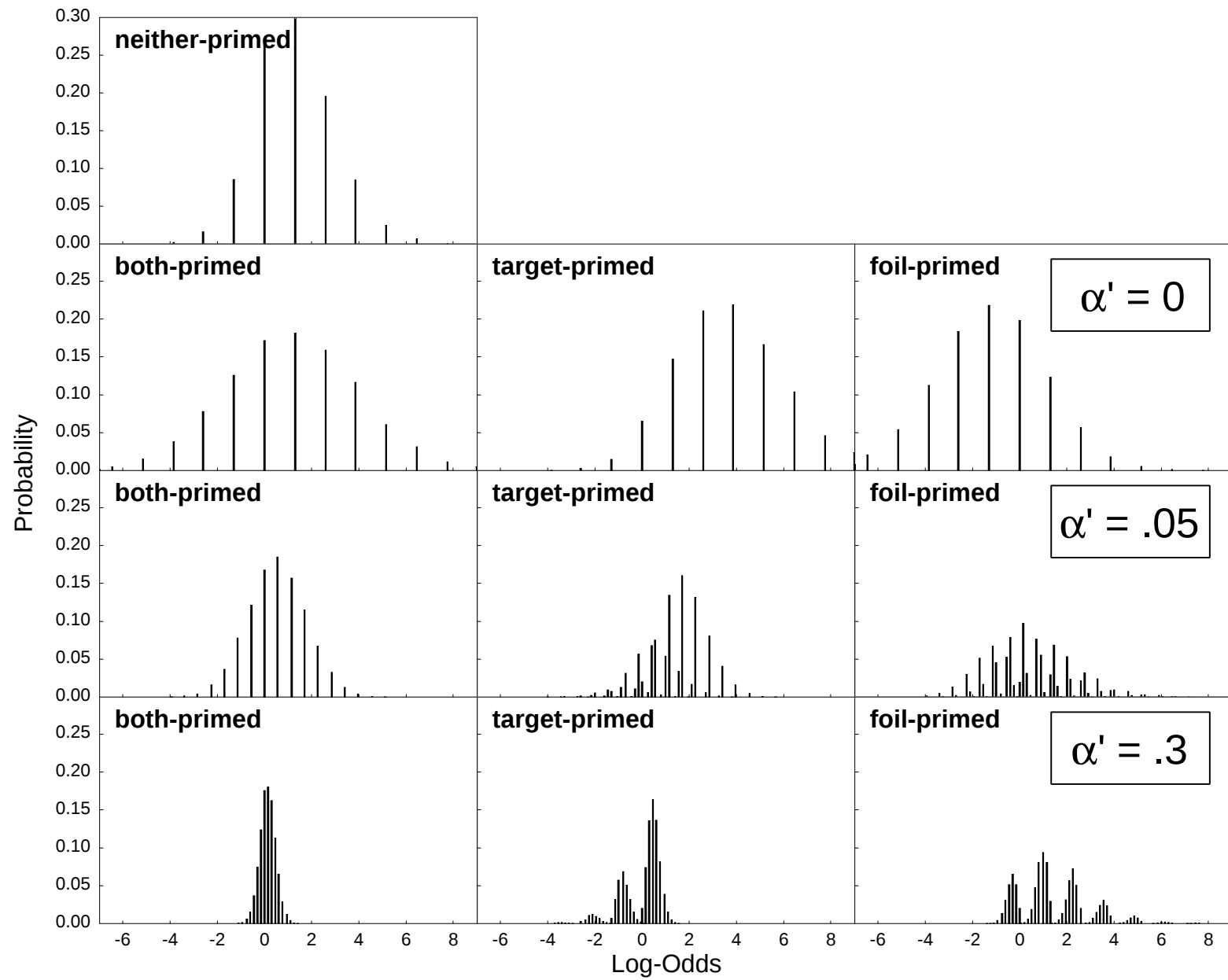


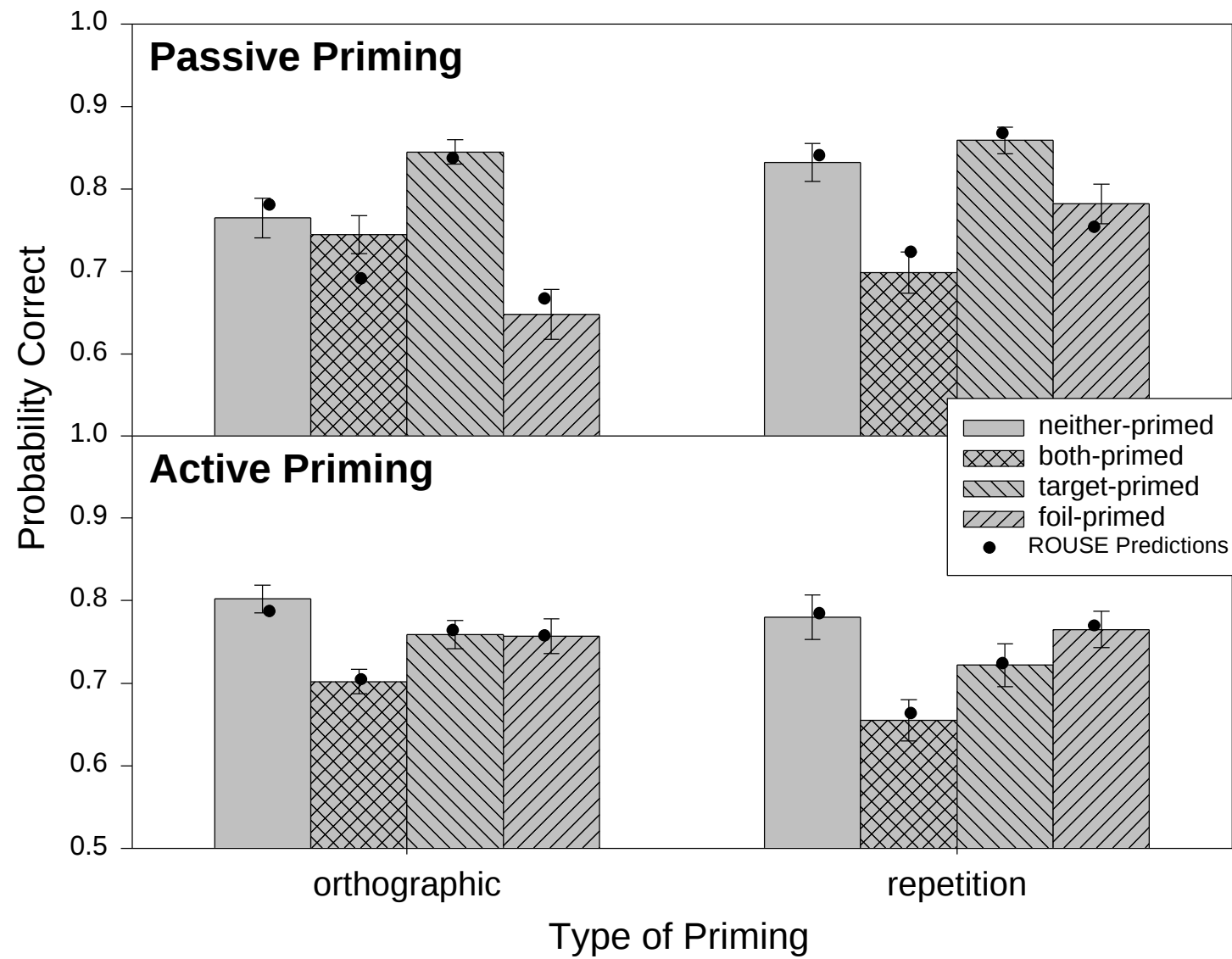


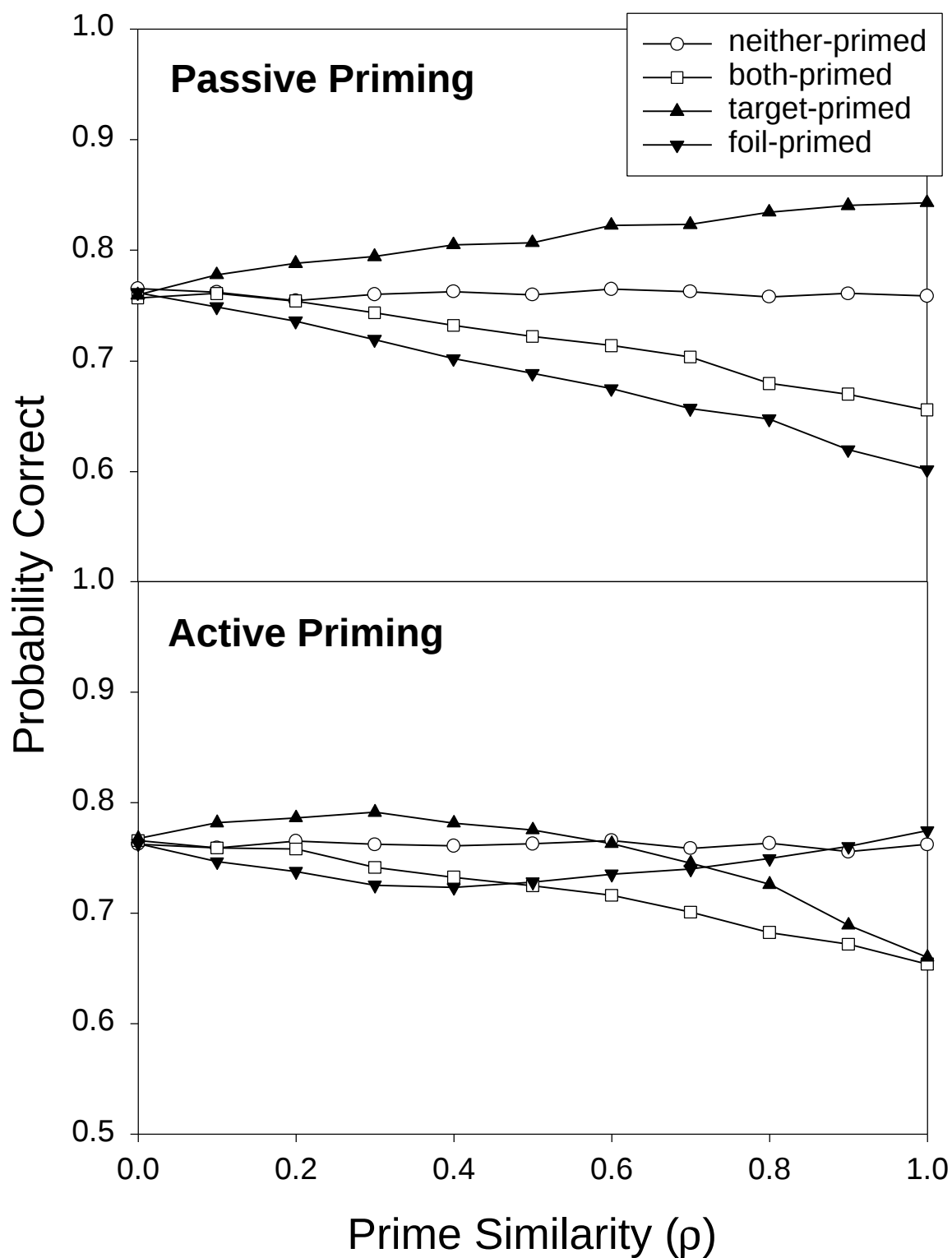


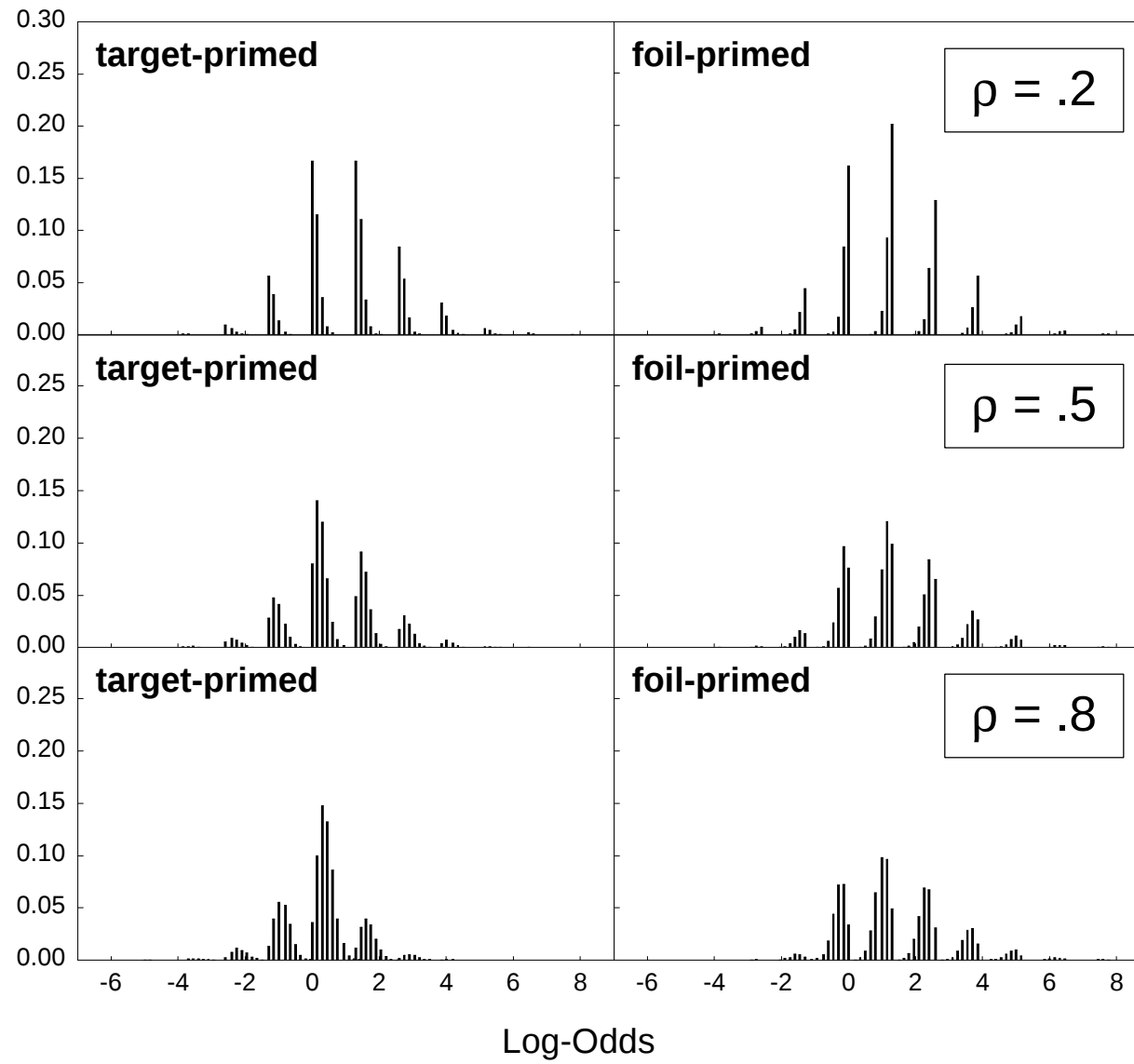
		State of Activation	
		OFF	ON
Appeared in Prime(s)	NO	$\frac{(1 - ?)(1 - \beta)}{(1 - ?)}$ $= (1 - \beta)$	$\frac{1 - (1 - ?)(1 - \beta)}{1 - (1 - ?)}$
	YES	$\frac{(1 - ?)(1 - a')(1 - \beta)}{(1 - ?)(1 - a')}$ $= (1 - \beta)$	$\frac{1 - (1 - ?)(1 - a')(1 - \beta)}{1 - (1 - ?)(1 - a')}$

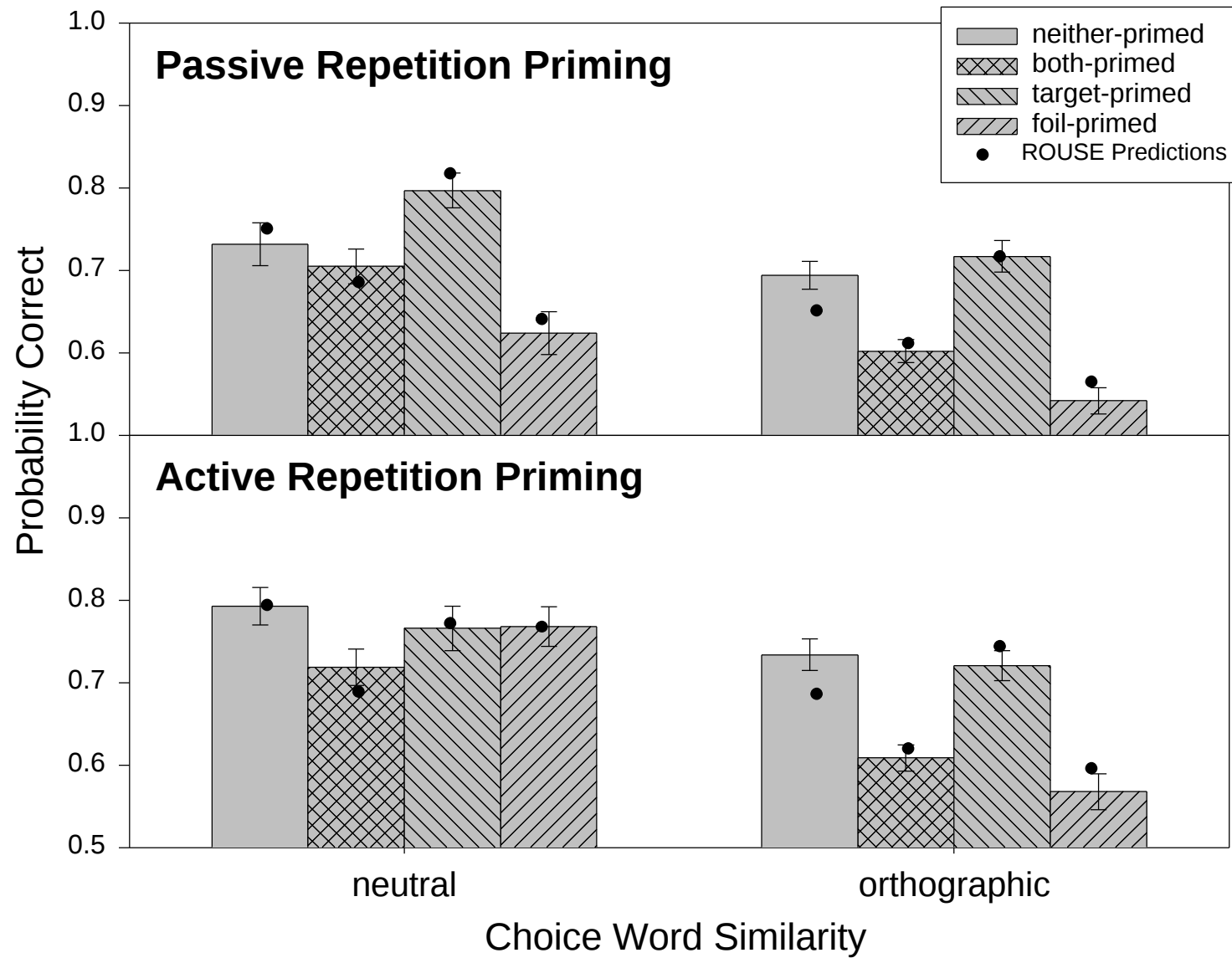


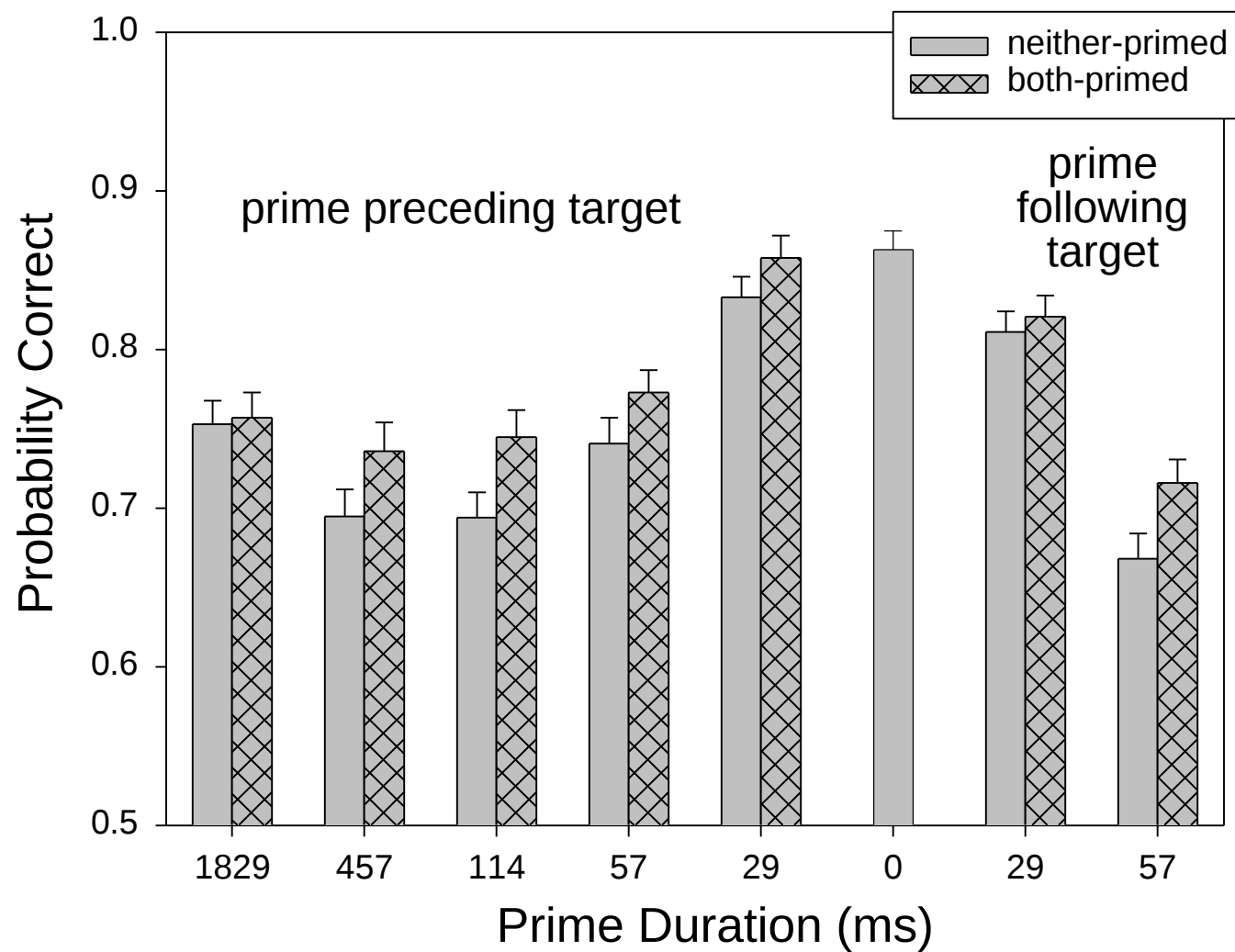


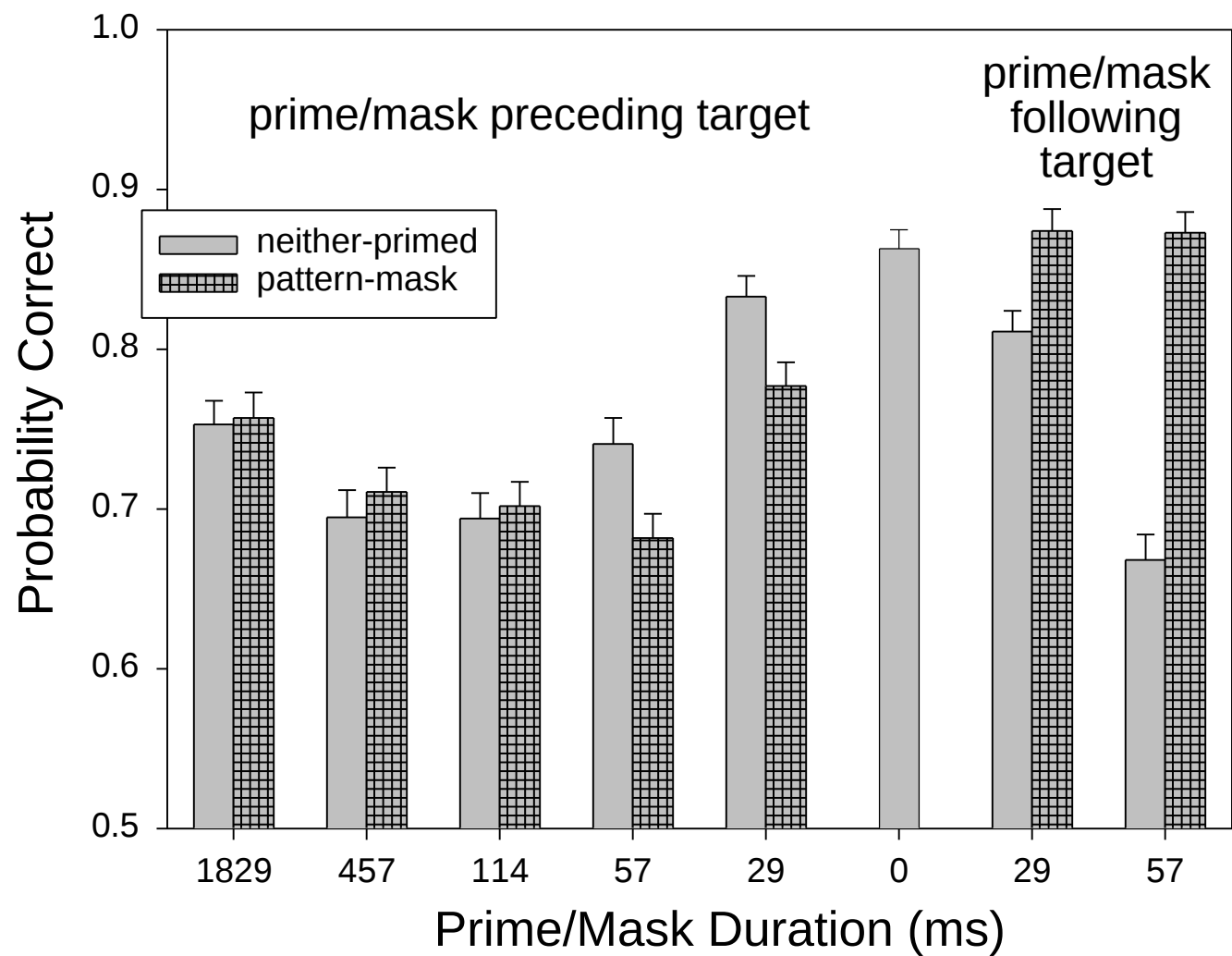


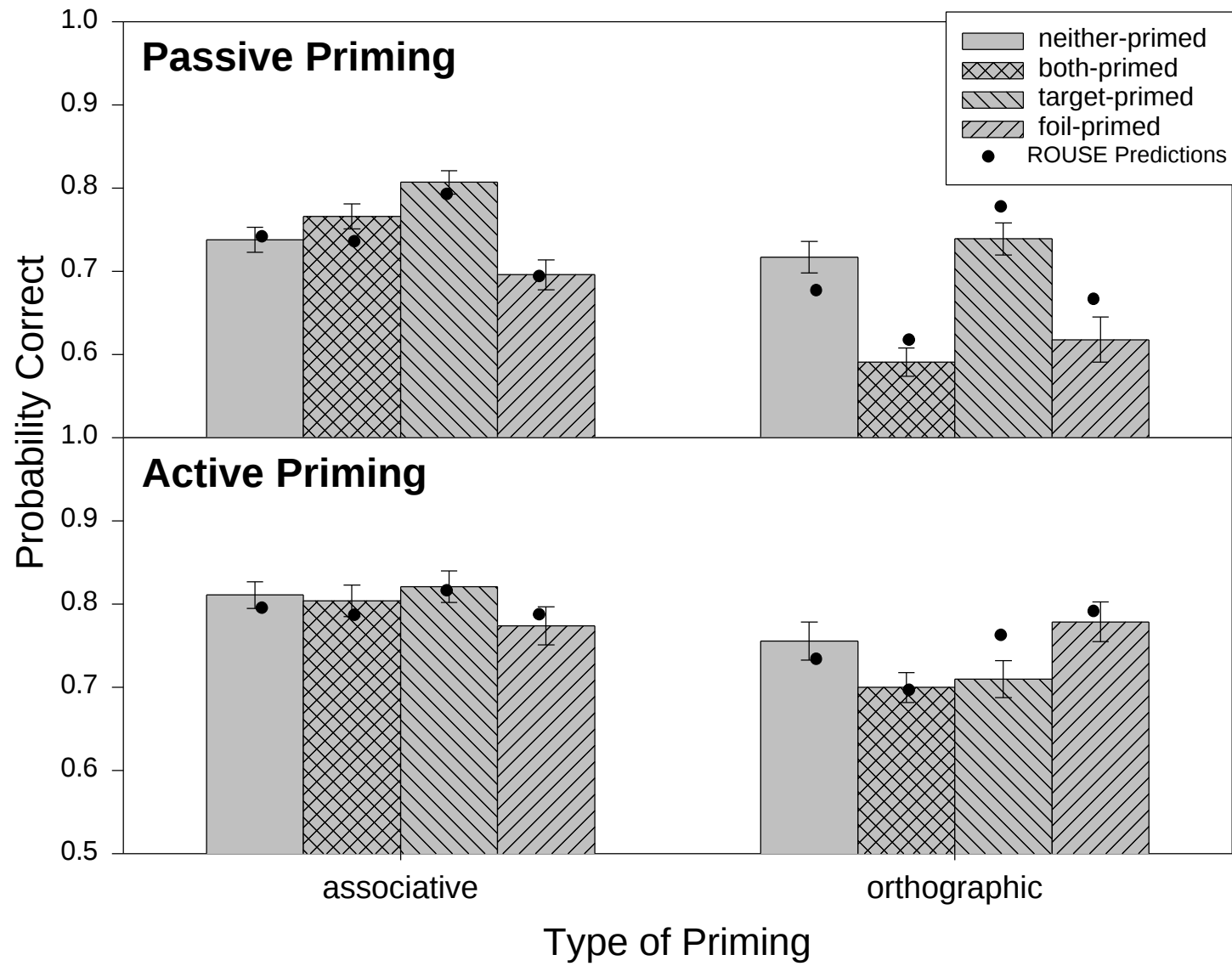


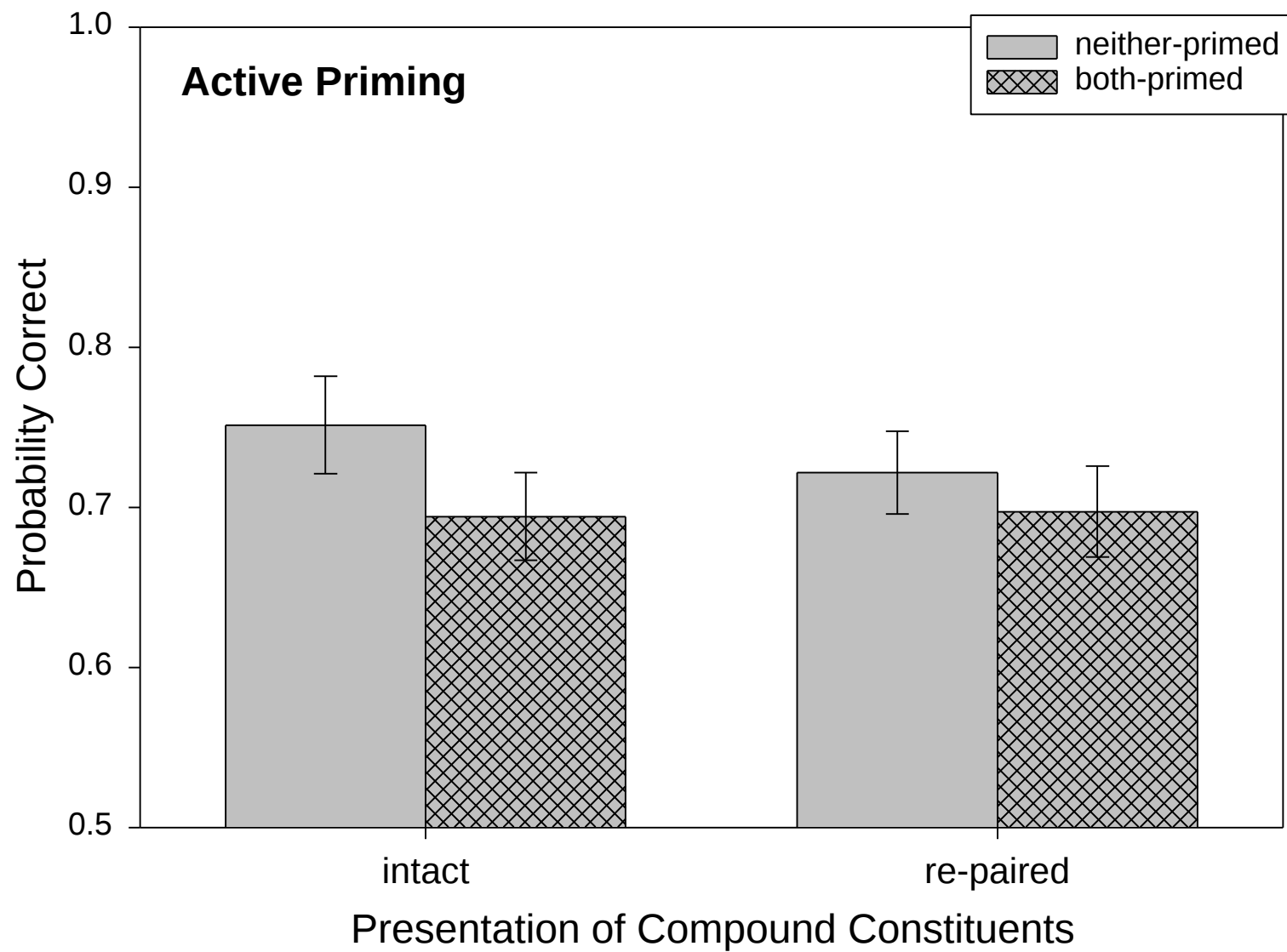


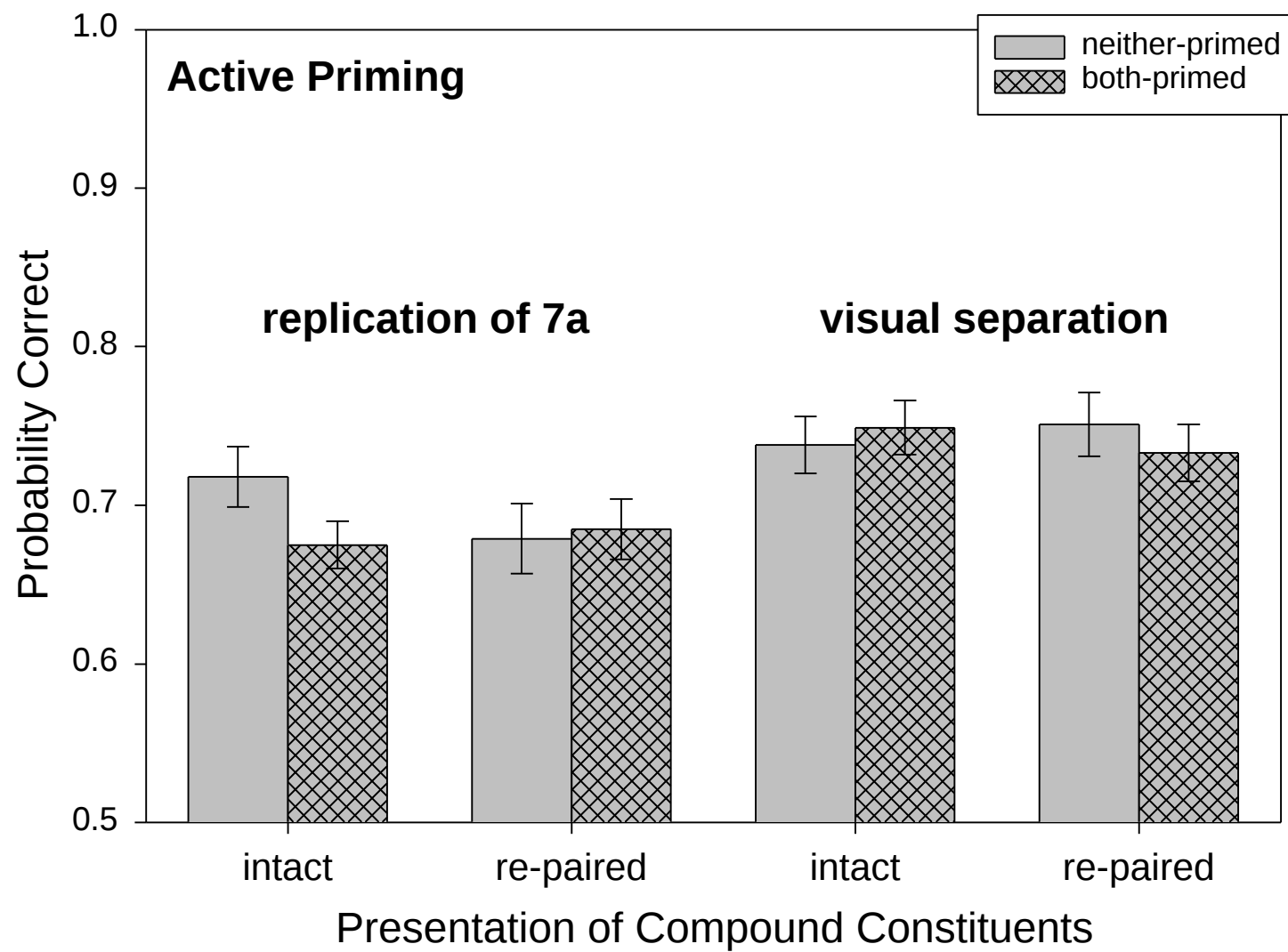












Curriculum Vitae

David E. Huber

<http://www.psych.indiana.edu/~dhuber/home.html>

University of Colorado, Boulder
Department of Psychology
Campus Box 345
Boulder, CO 80309-0345

dhuber@psych.colorado.edu
home: (303) 247-1759
lab: (303) 492-2269
fax: (303) 492-2967

Research Areas

- Implicit priming experiments distinguishing decisional and perceptual effects
- List length and strength effects in recognition memory
- Near-optimal Bayesian models of episodic and implicit memory
- Neural network and neural oscillation models of perception and memory

Education

- Ph.D. Joint degree in Cognitive Psychology and Cognitive Science with a minor in Neuroscience and a certificate in mathematical modeling and another in dynamical systems. Indiana University, 1994 - 1999. Advisor: Richard M. Shiffrin
- B.A. in Psychology and degree with Honors in Physics, Williams College, 1991.

Professional Employment

- Postdoctoral Fellow. University of Colorado, Institute of Cognitive Science, 1999-2001. Advisor: Randy O'Reilly
- Associate Instructor, "Methods of Experimental Psychology", 1998; 1996
- Research Assistant, Jeroen G. W. Raaijmakers, 1994
- Research Assistant, Richard M. Shiffrin, 1991-1994
- Research Assistant (part-time), Daniel B. Willingham, 1991

Honors and Fellowships

- University of Colorado, Institute of Cognitive Science, Postdoctoral Fellowship, 1999-2001
- Indiana University College of Arts and Sciences (COAS) Dissertation Year Research Fellowship for 1989 - 1999
- National Science Foundation (NSF) Graduate Research Fellowship (DGE-9253867 006), 1995 - 1998
- Workshop in Mathematical Psychology at UC Irvine (3 week selective workshop covering a variety of issues in mathematical psychology), 1997
- Commendation for second year of Ph.D. program, Indiana University, 1996
- Easter School on Computational Modeling at King's College (4 day selective workshop with a focus on memory), 1995
- National Institute of Health (NIH) pre-doctoral fellowship entitled "Modeling in Cognition" (T32 MA 19879-02), 1994-1995

Publications

- Huber, D. E. (1998). The development of synchrony between oscillating neurons. *Proceedings of the 20th Annual Conference of the Cognitive Science Society*.
- Shiffrin, R. M., Huber, D. E., & Marinelli, K. (1995). Effects of category length and strength on familiarity in recognition. *Journal of Experimental Psychology: Learning, Memory, & Cognition*, Vol. 21, No. 2, 267-287.
- Nobel, P. A. & Huber, D. E. (1993). Modeling forced-choice associative recognition through a hybrid of global recognition and cued-recall. *Proceedings of the 15th Annual Conference of the Cognitive Science Society*, 783-788.
- Huber, D. E., Marinelli, K., Ziemer, H. E., & Shiffrin, R. M. (1992). Does memory activation grow with list strength and/or length? *Proceedings of the 14th Annual Conference of the Cognitive Science Society*, 147-152.

Submitted Manuscripts

- Huber, D. E., Shiffrin, R. M., Lyle, K. B., & Ruys, K. I. (submitted). Perception and preference in short-term word priming. *Psychological Review*.

Manuscripts in Preparation

- Wagenmakers, E. M., Zeelenberg, R., Huber, D. E., Raaijmakers, J. G. W., Shiffrin, R. M., & Schooler, L. J. (in preparation). Bias vs sensitivity in short-term and long-term priming. In Marsolek, C. J. & Bowers, J. S. (Eds.), *Rethinking Implicit Memory*.
- Huber, D. E., Shiffrin, R. M., & Lyle, K. B. (in preparation). Active vs passive short-term priming: strategies and theories.

Conference Presentations

- Huber, D. E., Lyle, K. B., & Shiffrin, R. M. (1999). Short-term priming: data and a model for bias and interference *40th Annual meeting of the Psychonomic Society*, Los Angeles, California.
- Huber, D. E., Lyle, K. B., & Shiffrin, R. M. (1999). Short-term priming: data and a model for bias and interference. *32nd Annual Mathematical Psychology Meeting*, Santa Cruz, California.
- Huber, D. E. (1999). Short-term repetition and associative priming: bias or perception? *24th Annual Interdisciplinary Conference*, Teton Village, Jackson Hole, Wyoming.
- Huber, D. E. (1997). Entrainment as a model for visual adaptation and persistence. *30th Annual Mathematical Psychology Meeting*, Indiana University, Bloomington, Indiana.
- Shiffrin, R. M. & Huber, D. E. (1992). A Dynamic Model for Trace Activation, *25th Annual Mathematical Psychology Meeting*, Stanford University, Stanford, California.
- Shiffrin, R. M., & Marinelli, R. M. (1991) Interference and Forgetting in Recognition Memory, *33rd Annual meeting of the Psychonomic Society*, San Francisco, California. (note: Shiffrin, R. M. acknowledged Huber, D. E. in the presentation).

- Willingham, D. B., Huber, D. E., Spear, J. L. & Gabrieli, J. D. E. (1991). Mirror Tracing is Learned via a Series of Direction-specific Associations, *33rd Annual meeting of the Psychonomic Society*, San Francisco, California.

Technical Reports

- Huber, D. E., Shiffrin, R. M., Lyle, K. B., & Ruys, K. I. (1999). Perception and preference in short-term word priming. Technical Report #2??, Indiana University, Cognitive Science Program.
- Shiffrin, R. M., Huber, D. E., & Marinelli, K. (1993). Effects of Length and Strength on Familiarity in Recognition. Technical Report #94, Indiana University, Cognitive Science Program.

Software and Teaching Materials

- Diller, D., Huber D. E., Nobel, P., & Dennis, S. (1995). A MATRIX Model Simulator. *Noetica: Open Forum*, Vol. 1, Issue 6. <http://www.cs.indiana.edu/Noetica/OpenForumIssue6/>
- Huber D. E., Nobel, P. & Dennis, S. (1995). A TODAM Simulator. *Noetica: Open Forum*, Vol. 1, Issue 6. <http://www.cs.indiana.edu/Noetica/OpenForumIssue6/>

Professional Activities

Reviewing

- *Psychonomic Bulletin & Review*
- *Memory & Cognition*
- *Mathematical Social Sciences*

Teaching

- Methods of Experimental Psychology