

INHIBITION OF THE LITERAL: METAPHORS AND IDIOMS AS JUDGMENTAL PRIMES

ADAM D. GALINSKY

Northwestern University

SAM GLUCKSBERG

Princeton University

Four experiments demonstrate that priming effects depend on the context-appropriate meaning of the prime words. Most studies of semantic construct activation have presented prime words in contexts where the meaning of each word was invariant (e.g., word puzzles). In this research, we used words in contexts that supported either literal or figurative meanings, and found that only the context-appropriate meanings had subsequent priming effects on person-perception judgments. In Experiment 1, participants read the word “fire” in one of three contexts: a figurative use that implied recklessness (“playing with fire”), a figurative use referring to a hot streak (“on fire”), or a literal use (“playing by the fire”). Differential priming effects were obtained in a subsequent person-perception task that were consistent with the context-appropriate meanings of the priming expression. In Experiment 2, a conventional idiom, “break a leg,” produced divergent priming effects when used idiomatically than when used literally. In addition, we found evidence for inhibition of irrelevant literal meanings in Experiments 3a and 3b that provided support for the role of inhibitory processes in metaphor and idiom comprehension. Implications for how figurative language might differentially activate knowledge structures and the role of inhibitory processes in social perception are discussed.

A wide range of knowledge structures can be activated by situational contexts. Attitudes (Bargh, Chaiken, Gervernder, & Pratto, 1992), stereotypes (Devine, 1989), and even the self-concept (Bargh & Tota, 1988) have all been shown to be automatically activated by the mere presence

The authors thank Galen Bodenhausen and Thomas Mussweiler for their helpful comments on an earlier version of this article. The research was partially supported by Grant #SBR-9712601 from the National Science Foundation to Princeton University, S. Glucksberg, Principal Investigator.

Address correspondence to Adam Galinsky, Organizational Behavior, Kellogg School of Management, Northwestern University, Evanston, IL 60208-2001; E-mail: agalinsky@nwu.edu.

of a relevant object or symbol in the environment. Thinking processes can also be primed from constructs that suggest relationships among objects (Higgins & Chaires, 1980), to the mental simulation induced by exposure to a counterfactual event (Galinsky, Moskowitz, & Skurnik, in press). Activated knowledge structures made accessible through incidental exposure can often influence the interpretation of subsequent ambiguous events and behaviors. The pervasiveness of this phenomenon led Sedikides and Skowronski (1991) to conclude that there is a "fundamental law of cognitive structure activation." Given that the social world is often a buzzing and chaotic cauldron of complex information (Gilbert & Hixon, 1992; James, 1890) and that social perceivers prefer a state of epistemological certainty to a state of doubt and uncertainty (Heider, 1944; Moskowitz, Skurnik, & Galinsky, 1999), activated knowledge structures reduce doubt and simplify the complex world to manageable levels by providing individuals with an interpretive frame for understanding subsequent information (Stapel, Koomen, & van der Pligt, 1996).

Higgins, Rholes, and Jones (1977), in their seminal paper on knowledge-structure activation, had participants read words under the guise of an experiment in perception. In one condition, the words were related to recklessness, whereas in another condition they were related to adventurousness. The experimental participants were then asked to take part in a seemingly unrelated experiment on reading comprehension. The text that all participants read in this comprehension study was about a character named Donald, who was ambiguous on two opposing valenced trait constructs, recklessness and adventurousness. Those participants who had already been exposed to reckless-related words judged Donald to be more reckless than adventurous, whereas those participants who had seen the adventurousness-related words only judged Donald to be more adventurous than reckless. This priming effect was highly specific. In the Higgins et al. (1977) study, no effect was found when participants were primed with trait words that were either positive (e.g., satirical, grateful) or negative (e.g., listless, disrespectful) but were inapplicable to the Donald paragraph. Higgins et al. (1977) concluded that semantic primes exert an influence, not through the evaluative connotations of the traits, but through direct activation of a trait construct and its related meanings. "If a [primed] construct is applicable to the stimulus (i.e., there is sufficient match between the features of the construct and the features of the stimulus), then it will be used to encode or categorize the stimulus" (Higgins & Chaires, 1980, p. 351).

One criticism of the initial priming studies was that the knowledge structures were activated in artificial settings that were divorced from the typical contexts in which social information was encountered.

Moskowitz and Roman (1992) extended the ecological validity of both spontaneous trait inferences and previous priming studies by using stimuli that passively activated trait constructs during comprehension. They used trait-implying sentences that previous research (Uleman, Newman, & Moskowitz, 1996) had shown to produce spontaneous, unintentional trait inferences. Exposure to a single trait-implying sentence increased the accessibility of that trait construct, which was used to categorize nominally unrelated ambiguous behaviors.

Our research addresses the question of whether the same word would have differential priming effects depending on its context-determined meaning. Most previous priming studies have exposed participants to words divorced from any context, despite the ubiquity of context effects in all aspects of language comprehension (Austin, 1962; Anderson & Ortony, 1975; Gernsbacher, 1990; Glucksberg & Estes, *in press*; Simpson, 1984, 1994; Simpson & Krueger, 1991). Typical priming studies that have used words without supportive context include the Higgins et al. (1977) ostensible perception experiment or experiments by Herr (1986) and Moskowitz and Skurnik (1999), which embedded prime words in simple word puzzles. Figurative uses of words are an ideal route to explore the nature of context-determined priming effects because the meaning of a word used metaphorically often differs from, and may even oppose, its literal meaning. In addition, exploring the priming effects of figurative speech extends the work on the ecological validity of priming effects because metaphorical thinking is a ubiquitous part of mental life (Gibbs, 1994).

FIGURATIVE LANGUAGE

Metaphors pervade autobiographical accounts, particularly those with strong emotional content, and even novel metaphors occur more than once per minute during ordinary discourse (Gibbs, 1994). Metaphors (a) provide linguistic outlets for concepts that are difficult to convey using literal language; (b) are concise and satisfy the Gricean maxim of quantity (Grice, 1975); and (c) convey the complexity and vividness of phenomenological experience (Ortony, 1975). Gibbs (1994) has suggested that a limited number of metaphorical concepts help structure collective experience and provide a basis for conversational common ground. For example, the metaphorical concept "Anger is a heated fluid in a container" makes a number of idiomatic utterances comprehensible, varying from "I'm getting hot under the collar" to "He blew his top."

Despite the poetic nature of metaphors, their construction and comprehension do not appear to involve separate cognitive functions or systems than the production and understanding of literal utterances

(Gibbs, 1994; Glucksberg & Keysar, 1990). Glucksberg and Keysar (1990) have suggested that metaphors are class-inclusion assertions in which the metaphor vehicle refers to a diagnostic category constructed on the fly (Barsalou, 1983) and is a prototypical member of that category (Rosch, 1978). In the example "My job is a jail," the metaphor vehicle jail refers to a category lacking a conventional name but nonetheless shares a number of related properties, such as unpleasant, confining, unrewarding, difficult to escape, etc. Metaphorically used, jail refers to a type of thing rather than an actual physical entity. The properties of the jail constructed in the ad hoc category are then attributed to the metaphor target: job. It is the relevant properties of the metaphor vehicle and not the target that are constructed. When the vehicle and the target are reversed, either the result is anomalous (e.g., *My jail is a job*), or the meaning changes (Glucksberg, McGlone, & Manfredi, 1997). For example, the metaphor, "My surgeon is a butcher," ascribes the imprecision of the butcher's cutting style to the surgeon, thereby impugning the competence of the surgeon. When the metaphor is reversed to "My butcher is a surgeon," the butcher is being commended for the deftness of his or her cutting. The category that is constructed during metaphor comprehension should serve a substantively similar role to the activation of stored knowledge structures in providing an interpretative frame for comprehending subsequently encountered social information.

Given the pervasive and unintentional nature of judgmental and trait priming, one would expect figurative primes to be as effective as literal ones. But there is a potential problem with figurative-language primes. By their very nature, metaphors and idioms are inherently ambiguous: all figurative expressions also have contextually inappropriate literal meanings. Would both the literal and figurative meanings affect subsequent trait judgments, or will such judgments be affected only by the context-appropriate meaning of the expression? This is an especially important question when the potential semantic constructs activated during the literal and figurative use of a phrase are in opposition to one another.

Consider the standard Donald paradigm (Higgins et al., 1977). Donald is ambiguous on the two opposing traits, adventurous and reckless. Would an expression that used the same critical word produce opposite priming effects when that word is used in two different metaphors (i.e., one that referred to dangerous and risky behaviors, the other that referred to adventurous behavior)? Experiment 1 examined the effects of using the word fire in three different contexts. One context was the metaphor "You're playing with fire," which implies that the addressee is engaged in risky behavior. A second context was the metaphor "You're on fire," used to refer to someone who is on a winning streak and can do no wrong. A third, literal context was included to provide a neutral

baseline against which to assess the differential effects of metaphors as judgmental primes.

EXPERIMENT 1: ARE METAPHOR PRIMES CONTEXT-SPECIFIC?

Participants were exposed first to a story that concluded either with a literal reference or with one of two opposing metaphors with respect to the dimension of reckless-adventurous. Participants then completed a person perception task in which they formed an impression of a target individual who was ambiguous for the trait dimension of reckless-adventurous. If metaphors provide differential property activation that is context-appropriate, then participants should rate this ambiguous target person as more reckless after exposure to a figurative expression that implied recklessness (Playing with fire), but more adventurous after exposure to a figurative expression that implied invulnerability from harm (You're on fire) relative to a literal control condition (He sat by the fire).

METHOD

PARTICIPANTS AND DESIGN

Participants were 67 undergraduates at Princeton University who received credit as part of a course requirement. All were native English speakers. The study used a single-factor, between-subjects design with three levels of the manipulated variable (literal reference to fire, "Sitting by the fire"; figurative use of fire implying danger, "You're playing with fire"; figurative use of fire implying winning streak, "You're on fire"). The prime scenarios and the Donald impression-formation story are presented in the Appendix.

PROCEDURE AND STIMULUS MATERIALS

Upon arriving at the laboratory, participants were told that the experiment was concerned with how a delay might affect individuals' reactions to stories and events.

First, participants were presented with the prime stories. They were told to read the story carefully because they would be asked questions about it later. Next, participants were given a counting backwards task in order to clear working memory (Higgins et al., 1977). The two pages after the counting backwards task presented the impression formation task. Previous research has found that when the priming and judgment tasks are not sufficiently distinct, participants attempt to control for any

effects that the prime might have on their subsequent judgments (Martin, 1986). To minimize this possibility, the two pages of the impression formation task were presented in a different font from the prime scenario (Galinsky et al., in press). We anticipated that the differing fonts would foster the perception that the two tasks were unrelated. Participants were instructed to form an impression of Donald as they read the paragraph. On the following page, participants were asked to characterize Donald on a 9-point scale. The rating scale was anchored at (1) adventurous and (9) reckless. After rating Donald, participants were asked a few questions about the central character from the prime scenarios. Participants were asked what type of activity the main character from the prime scenario was practicing for and what type of cards the characters were playing. After completing the experiment, participants were queried verbally about their suspicions. None of the participants mentioned any connection between the prime scenarios and the Donald paragraph. When asked if their judgments of Donald could have been influenced by the prime scenarios, none of the participants correctly identified the hypothesized mechanisms.

RESULTS AND DISCUSSION

Judgments of Donald differed as a function of prime type, $F(2, 64) = 6.5, p < .01$. Donald was rated the most reckless ($M = 5.9$) after the prime "you're playing with fire" and the least reckless ($M = 4.0$) after the prime "You're on fire." Judgments of Donald after the control prime in which a literal reference to fire was made fell between the two metaphorical primes ($M = 4.8$). Primed with a metaphor that activated the construct dangerousness, participants judged Donald to be reckless. Previous exposure to figurative speech that implied immunity from harm led to judgments of Donald as less reckless and more adventurous. A linear contrast was significant, $F(1, 64) = 13.00, p < .001$, and the residual was not, $F(1, 64) < 1$. The a priori hypothesis is thus strongly supported and is sufficient to account for the nonchance variation in the data (Abelson & Prentice, 1998).

The fact that metaphors can serve as differential, context-appropriate primes suggests that the priming consequences of exposure to a word are not invariant but depend on its contextually appropriate meaning. A metaphor that implied danger resulted in judgments of Donald as relatively more reckless, whereas metaphors that implied immunity from harm resulted in judgments that Donald was less reckless. The differential effects of the primes suggest that context-appropriate, metaphor-relevant properties were activated during metaphor comprehension.

EXPERIMENT 2: ARE IDIOM PRIMES CONTEXT SPECIFIC?

Experiment 2 was designed to explore whether the priming results were confined to metaphors or generalized to other forms of figurative speech such as idioms. Idioms can be considered to be dead metaphors and frozen, formulaic phrases directly represented in the lexicon (Bobrow & Bell, 1973), understood by retrieving from memory the meaning of the idiom as a whole. This is especially true of idioms that are noncompositional (i.e., ones that can not be broken down into its constituent meanings). For example, "Kick the bucket" is known to refer to dying even though none of its elements have any connection to dying. Idioms are often understood immediately: McGlone, Glucksberg, and Cacciari (1994) found that participants comprehended idioms more quickly than either their literal paraphrases or their variants. This implies that idioms are comprehended by directly retrieving its meaning from memory without requiring linguistic processing. Whereas properties of metaphor vehicles are attributed to metaphor topics, idioms are self-contained lexical units that do not depend on property attribution.

To assess whether idioms could also serve as primes, we used the expression "Break a leg" in two contexts. In the figurative-appropriate context, "Break a leg" was used as an expression of good luck and an exhortation to do well. In a literal-appropriate context, breaking a leg was used to refer to someone who suffers a broken leg. We expected people to rate Donald, who "has risked injury and death many times," as more reckless after a scenario in which someone actually breaks a leg (Galinsky et al., in press) than after a scenario in which someone uses the expression "Break a leg" as a way of expressing good luck.

METHOD

PARTICIPANTS AND DESIGN

Participants were 42 undergraduates at Princeton University who received credit as part of a course requirement. All were native English speakers. The study used a single-factor, between-subjects design with three levels of the manipulated variable (literal reference to a broken leg; idiomatic expression "Break a leg" used to assert sentiments of good luck; and a no-prime control).

PROCEDURE AND STIMULUS MATERIALS

Upon arriving at the laboratory, participants were given the same instructions as participants in Experiment 1. They were presented with one of two prime scenarios or with no scenario at all (a no-prime con-

trol). The literal reference scenario described a man breaking his leg while rushing off to a performance. In the scenario that used the idiom, the man's roommate said, "Break a leg," as he rushed off to an unspecified performance. Participants were again instructed to form an impression of Donald as they read the paragraph and to rate Donald on a scale that was anchored at (1) adventurous and (9) reckless. As in Experiment 1, none of the participants noticed any connection between the priming paragraph and the impression formation paragraphs when queried about their suspicions after the experiment had concluded. The prime materials are presented in the appendix.

RESULTS AND DISCUSSION

Judgments of Donald differed as a function of prime type, $F(2,39) = 3.4, p < .05$. Donald was rated as most reckless after exposure to a literal reference to a broken leg ($M = 5.8$) than following the idiomatic reference ($M = 4.0$) and the no prime control ($M = 4.6$). A linear contrast was significant, $F(1, 39) = 6.7, p < .02$ and the residual was not, $F(1,39) < 1$. Literal reference to a broken leg increased ratings of Donald's recklessness. In contrast, the idiomatic use of "break a leg" unexpectedly appeared to inhibit the literal meaning, producing decreased judgments of recklessness relative to the control condition.

Why would the figurative usage inhibit the literal meaning of the figurative phrase? When figurative language is understood, either in ordinary conversation or in psycholinguistic experiments, irrelevant meanings appear to be efficiently filtered. In ordinary conversation, when someone says, "My lawyer is a shark," it is understood that the lawyer may be aggressive, predatory, and tenacious, but no one would take the expression to also mean that the lawyer has fins and can breathe under water. In experimental contexts, people seem to be able to inhibit irrelevant literal information while activating and retaining relevant metaphorical information. Gernsbacher, Keysar and Robertson (1995; see Glucksberg, Newsome, & Goldvarg, 1997) provide a convincing demonstration of the differential accessibility of metaphor-relevant and irrelevant information during metaphor comprehension. Participants in their experiment read sentences, one at a time, and judged whether each sentence was sensible. After reading a metaphor such as "My lawyer was a shark," participants comprehended metaphor-relevant property sentences (e.g., "sharks are vicious") more quickly than metaphor-irrelevant ones (e.g., "sharks are good swimmers"). The reverse was true after reading literal sentences about sharks, such as "The hammerhead is a shark." Gernsbacher et al. (1995) concluded that metaphor comprehension involves the simultaneous accessibility of metaphor-relevant information and the active inhibition of meta-

phor-irrelevant information. The results of the first two experiments suggest that the differential accessibility of metaphor and idiom-relevant and metaphor and idiom-irrelevant information influences unrelated judgments.

EXPERIMENTS 3A AND 3B: ARE FIGURATIVE-IRRELEVANT PROPERTIES INHIBITED?

The results of Experiment 2 suggest that the active inhibition of figurative-irrelevant properties also function when idioms are used as primes. But, there is an alternative explanation for this result. It may well be the case that the expression "Break a leg" as a way of wishing someone good luck does not inhibit recklessness but simply implies adventurousness, similar to a straightforward and literal expression such as good luck. If so, then wishing Donald well by saying, "Break a leg," should increase participants' judgments of adventurous. Experiment 3a was designed to discriminate between these two alternatives: an inhibitory effect of the "Break a leg" expression that reduces judgments of recklessness versus a positive priming effect of this expression on judgments of adventurousness. If "Good luck" has the same effect as Break a leg on judgments of recklessness-adventurousness, then the break a leg effect is most likely because of positive priming of adventurousness. Conversely, if the good luck condition is equivalent to a no-prime control, then any effects of break a leg on judgements of reduced recklessness would be best attributed to negative, inhibitory priming rather than positive priming of adventurousness. Experiment 3a was a replication of Experiment 2 with an additional control condition. The literal expression "Good luck" contrasted with the figurative expression "Break a leg," and a no-prime control.

METHOD

PARTICIPANTS AND DESIGN

Participants were 68 students recruited from a university library. All were native English speakers. The study used a single-factor, between-subjects design with three levels of the manipulated variable (idiomatic expression "Break a leg," literal expression "Good luck;" and a no-prime control).

PROCEDURE AND STIMULUS MATERIALS

Participants were given the same written instructions that participants had received in the previous experiments. Two groups of participants

were presented with one of two prime scenarios. In the scenario that used the idiom, the man's roommate said, "Break a leg," as he rushed off to his performance. The other prime scenario had the roommate saying, "Good luck," as he rushed from the room. A third group of participants was given the person perception task without any prime scenario. Participants were again instructed to form an impression of Donald as they read the paragraph and to rate Donald on a scale that was anchored at (1) adventurous and (9) reckless. As in the previous experiments, none of the participants noticed a connection between the primes and the person perception task.

RESULTS AND DISCUSSION

Judgments of Donald differed marginally according to type of prime, $F(2,65) = 2.1, p = .13$. Participants primed with the idiom break a leg, rated Donald as less reckless ($M = 4.4$) than participants primed with good luck ($M = 5.4$) or participants in the no-prime control condition ($M = 5.3$). A comparison of the idiom prime condition against the other two conditions was significant, $F(1, 65) = 4.2, p < .05$ and the residual was not, $F(1, 65) < 1$. These results suggest that the idiomatic prime produced an inhibition effect, making the construct less accessible and therefore less likely to be used to categorize subsequently encountered ambiguous behaviors applicable to that construct.

But there is still an alternative explanation for these results. Because participants used a single-rating scale, any decreased judgment of recklessness can be interpreted as an increased judgment of adventurousness and vice-versa. Thus, our results can be interpreted as a positive priming effect of "break a leg" on adventurousness ratings instead of a negative inhibitory effect on recklessness ratings. To assess whether judgments of recklessness are inhibited or judgments of adventurousness are enhanced by the idiom break a leg, Experiment 3b provided independent judgments of recklessness and adventurous by using two separate scales: (a) A rating scale for adventurousness for one group of participants; and (b) a rating scale for recklessness for a second group of participants.

METHOD

PARTICIPANTS AND DESIGN

Participants were 86 students recruited from a university library. All were native English speakers. The study used a 3(prime: idiomatic/literal/no prime control) $\times$ 2(response scale: adventurous/reckless) factorial design.

PROCEDURE AND STIMULUS MATERIALS

Participants were given the same written instructions that participants had received in the previous experiments. Two groups of participants were presented with one of two prime scenarios. In the scenario that used the idiom, the man's roommate said, "Break a leg," as he rushed off to his performance. The other prime scenario had the protagonist break his leg as he rushed from the room. A third group was given the person perception task without any prime scenario. Participants were instructed to form an impression of Donald as they read the paragraph. Half of the participants was asked to rate Donald on a scale that was anchored at (1) not at all adventurous and (9) very adventurous. The other half of the participants was asked to rate Donald on a scale that was anchored at (1) not at all reckless and (9) very reckless. As in the previous experiments, none of participants noticed a connection between the primes and the person perception task.

RESULTS AND DISCUSSION

A 3 (prime: idiomatic/literal/no-prime control) $\times$ 2 (response scale: adventurous/reckless) between-participants ANOVA was run on impressions of Donald. A main effect for prime, $F(2,80) = 7.2, p < .01$, was qualified by a significant prime $\times$ scale interaction, $F(2,80) = 8.7, p < .01$ (see Figure 1). One-way ANOVAs were conducted for each response scale. For the recklessness scale, there was a significant effect of prime, $F(2,40) = 14.6, p < .01$. The literal use of break a leg as a prime led to increased judgements of Donald's recklessness ($M = 8.5$) compared with the no-prime control ($M = 7.5$), $t(40) = 2.9, p < .01$. In addition, the idiomatic prime led to decreased judgements of Donald's recklessness ($M = 6.6$), compared with the no-prime control, $t(40) = 2.3, p < .03$. For the adventurous scale, there was no effect of prime on judgments of Donald, $F(2,40) = 1.3, p < .29$. These results suggest that the idiomatic prime produced an inhibition effect, making the construct less accessible and therefore less likely to be used to categorize subsequently encountered ambiguous behaviors applicable to that construct.

GENERAL DISCUSSION

The experiments reported here suggest that it is only the context-appropriate meanings of words that are activated in priming experiments. Thus, the judgmental consequences of exposure to any particular word are not invariant but depend on context-based usage of



FIGURE 1. Ratings of recklessness and adventurousness as a function of prime type.

the word. The experiments reported here have implications in two domains: language processing and judgmental priming.

With respect to metaphor and idiom processing, our results are consistent with language comprehension models that posit differential activation of discourse-relevant versus discourse-irrelevant information. According to Gernsbacher's (1990) structure-building model of language comprehension for example, relevant material is enhanced and irrelevant material is actively inhibited in order to generate coherent discourse representations (see also Gernsbacher & Faust, 1990; Hasher & Zacks, 1988; Simpson & Kang, 1994). Gernsbacher et al. (1995) found that both enhancement and inhibition are involved in metaphor comprehension. They concluded that metaphor-irrelevant properties are not merely not activated, but may be actively inhibited during metaphor comprehension (see Kintsch, 1998, for an analogous mechanism from a computational viewpoint).

In a follow-up study, Glucksberg et al. (1997) replicated Gernsbacher et al. (1995) but with an important difference in their materials. Glucksberg et al. (1997) adapted the independent cue paradigm, a procedure devised by Anderson and Spellman (1995) to directly assess inhibitory effects. This technique effectively excludes alternative explanations of negative priming effects, such as task-specific memory strategies. As in the Gernsbacher et al. (1995) experiment, participants read metaphors and literal control sentences, and then judged whether

metaphor-relevant and irrelevant probe sentences made sense. Unlike the original experiment, the probe sentences did not use the priming sentence predicate. For example, if a metaphor prime had been "My lawyer is a shark," the metaphor-irrelevant probe sentence was "Geese are good swimmers" instead of "Sharks are good swimmers." Because none of the words of the prime sentence are repeated in the probe sentence, experimental participants should not be cued to use information generated while reading the prime sentence. Instead, the independent-cue probes should directly tap relative levels of activation of the metaphor-relevant and -irrelevant property concepts. Thus, if the property "swimming" had been actively inhibited during metaphor comprehension, then the response to the geese swimming assertion should be slower when it follows a metaphor than when it follows a literal statement such as "The hammerhead is shark." Similarly, metaphor-relevant probe sentences, such as "Geese can be tenacious," can be used to determine if metaphor-relevant properties are actively enhanced (Gernsbacher, 1990). As expected, metaphor-irrelevant properties were selectively inhibited. For example, in the metaphor "My lawyer is a shark," properties that were only associated with literal sharks and are thus metaphor-irrelevant (e.g., the skilled swimming of sharks) were actively inhibited. Our impression formation task can be construed as an analog of the independent cue technique. There is no obvious connection between the Donald story, which is analogous to the probe sentences in the independent cue paradigm, and the protagonist in the earlier priming scenarios, yet impressions of Donald were influenced by the relevant material in those scenarios.

Taken as a whole, our current findings of positive and negative priming effects provide additional evidence that metaphors and idioms produce enhancement of relevant and inhibition of irrelevant information. Furthermore, this differential activation can affect subsequent evaluations and impressions of apparently unrelated objects.

With respect to judgmental priming, our findings argue for the ecological validity of priming effects (Moskowitz & Roman, 1992). Metaphors and idioms are frequently used in everyday discourse (Gibbs, 1994), and so constructs activated during metaphor comprehension could affect a wide variety of subsequent judgments and behaviors. The more frequently a construct is activated the stronger and more lasting are its accessibility effects (Srull & Wyer, 1980). Often, the same concept is expressed and elaborated several times in discourse through different instantiations of the same conceptual metaphor. Allbritton, McKoon, and Gerrig (1995) demonstrated some cognitive effects of activating a particular metaphorical structure. After reading a scenario

that contained several instantiations of a metaphorical schema, participants were quicker to recognize subsequent metaphor-related words and sentences than after reading comparable scenarios that did not instantiate any metaphors. One such scenario referred to crime as a disease, containing metaphors such as "The city's crime epidemic was raging out of control" and "...the violence began to infect even safe neighborhoods" (Albritton et al., 1995, p. 613). After reading this scenario, responses on a delayed recognition test were faster to the target "Public officials desperately looked for a cure." It may well be that describing a topic such as crime in terms of disease could activate the disease concept sufficiently to influence subsequent judgments in the same way that judgments of Donald were influenced. For example, after talking about crime as a disease, would people then tend to frame such topics as obesity, envy, or any other human foible as a disease? Our results, combined with those of Albritton et al. (1995), suggest that this would be a real possibility. If this does happen, then we would have an ironic combination of an adaptive and efficient comprehension mechanism leading to nonconscious, biased evaluation of subsequent, unrelated objects of attention.

PRIMING AND INHIBITION

These results extend early work on priming by demonstrating that the priming consequences of any given word depend critically on the context in which it is used, and extend recent work exploring the role of inhibitory mechanisms in everyday social perception. Macrae, Bodenhausen, and Milne (1995) found that when perceivers encountered a target that could be categorized with reference to multiple social categories, the context led to the facilitation of one construct and the active inhibition of the other potential construct. They found that after participants saw a videotape of a Chinese woman eating noodles with chopsticks (i.e., the contextualized behavior was relevant to the Chinese stereotype) they responded more quickly to stereotypic words related to Chinese people, whereas participants responded more slowly to stereotypic words related to women, relative to control words that were unrelated to the either Chinese or women stereotypes. These inhibition and facilitation processes appeared to operate in parallel.

Similarly, Dijksterhuis and van Knippenberg (1996) found that priming a stereotype facilitated access to stereotype-consistent traits and inhibited access to stereotype-inconsistent traits. Primed with the stereotype of soccer hooligans, participants displayed increased access to words such as aggressive and decreased access to words such as in-

telligent relative to both a no-prime control condition and irrelevant words. The results of the experiments presented here are conceptual replications of the Dijksterhuis and van Knippenberg (1996) results. Just as "intelligent" is inconsistent with the stereotype of soccer hooligans, the figurative meanings of "break a leg" and "on fire" are inconsistent with their literal meanings. Rothbart, Sriram, and Davis-Stitt (1996) obtained similar results for category exemplars; category label primes increased the ability to retrieve typical category members, assessed through reaction time and error rates while limiting access to atypical members. The more counterstereotypic a category member was perceived to be, the lower the probability that this member would be activated and retrieved when primed with the stereotype label. The role of typicality in exemplar retrieval was explored further in a field study involving rival fraternities; typical members of the rival fraternities were more likely to be recalled, even when controlling for familiarity and liking. Like the Rothbart et al. (1996) field experiment that demonstrated the everyday consequence of inhibitory mechanisms with regard to stereotypicality, our results show that the inhibitory processes endemic to language processing can result in altered perceptions of subsequently encountered social stimuli.

The majority of previous research exploring the role of inhibition in social perception has relied on reaction time measures that allow for precision but do not speak to the applicability or consequences of the effects. By replicating the results in the person perception paradigm, we have established that the accessibility and inhibition of constructs during metaphor comprehension can have real judgmental consequences and can alter social perception with the possibility of affecting interpersonal interactions and social behavior (Bargh, Chen, & Burrows, 1996; Chen & Bargh, 1997).

CONTRAST EFFECTS, ALTERNATIVE EXPLANATIONS AND FUTURE STUDIES

We have claimed that the effects of decreased recklessness after figurative primes are the direct result of inhibitory processes involved in figurative language comprehension. The pattern of results is also consistent with contrast effects, the tendency to rate subsequent targets as less, rather than more, similar to activated knowledge structures. Contrast effects tend to occur in one of two situations (Moskowitz & Skurnik, 1999; Stapel, Koomen, & van der Pligt, 1997): (a) when participants recognize the potential biasing effect of the prime and attempt to correct for this influence (Martin, 1986); or (b) when the primes serve as a stan-

dard of comparison (Herr, 1986). Correction-produced contrast occurs when the priming episode involves trait terms and is particularly blatant (i.e., when the prime words have clear connection to the judgment task) (Strack, Schwarz, Bless, Kübler, & Wanke, 1993; Wegener & Petty, 1995). Standard of comparison-induced contrast occurs after exposure to exemplars; for example, ambiguously hostile behaviors encountered during the judgment task are judged to be less hostile after previous exposure to extremely hostile exemplars such as Adolf Hitler or Charles Manson.

The descriptions of the processes involved in contrast effects (see Moskowitz and Skurnik, 1999) imply that contrast effects are generated at the time of the judgment and not at the time of the primes. In support of the notion that contrast effects occur at the time of judgment, Moskowitz and Skurnik (1999) eliminated contrast effects by placing participants under cognitive load at the judgment stage. We have suggested that our effects are driven by inhibition effects that occur at the time of primes: the literal constructs receive decreased levels of activation below baseline levels. Future research should explore whether the inhibition effects are impervious to divided attention during the judgment stage, which would suggest that the effects are driven by inhibitory processes during the priming episode and not by contrast effects at the time of judgment. In addition, the results should be replicated with different trait dimensions and person perception materials to assess the generalizability of the results. Finally, future studies should also use reaction time measures to establish more firmly that the literal meanings were actively inhibited, and such results would also serve to differentiate the role of inhibition and contrast effects in the judgmental consequences of figurative primes. The reaction time measure that Macrae et al. (1995) used suggests that the activating one social stereotype inhibited a competing social stereotype at the time of activation.

CONCLUSION

The research reported in this article demonstrates that figurative language can serve as primes and that the priming effects of words are not invariant but depend on the context-appropriate meaning of the word. Our results also suggest that understanding inhibitory processes is necessary for a complete understanding of social perception.

APPENDIX: EXPERIMENTAL MATERIALS

1. *Donald Paragraph*

Donald spent a great amount of his time in search of what he liked to call excitement. He had already climbed Mount McKinley, driven in a demolition derby, shot the Colorado rapids in a kayak, and piloted a jet-powered boat—without knowing very much about boats. He had risked injury, and even death, a number of times. Now he was in search of new excitement. Other than business engagements, Donald's contacts with people were rather limited. He felt he didn't really need to rely on anyone.

2. *Prime Scenario for Experiment 1 with Metaphor Use of Fire Implying Winning Streak*

As the snow finally began to die down, John went through his practice routine one last time. He had gotten to the point where his nervousness had become excitement. About one hour before he was to leave, he shared a cup of tea with his roommate. They sat in the living room with the big bay window watching the snow begin to fall more heavily again. They decided to play a game of cards, because they were his favorite game. While playing, his roommate finally broached a sensitive subject about John's interest in their friend's girlfriend. On this particular night, John was unstoppable, winning every game. His roommate said, "Man, you're on fire."

3. *Ending of Prime Scenario for Experiment 1 with Metaphor Use of Fire Implying Recklessness*

On this particular night, John was unstoppable, winning every game. While playing, his roommate finally broached a sensitive subject about John's interest in their friend's girlfriend. His roommate explained to John, "Man, you're playing with fire."

4. *Ending of Prime Scenario for Experiment 1 with "Sitting by the Fire Control"*

They decided to play a game by the fire; cards were his favorite game.

5. Prime Scenario for Experiment 2 with Idiom

As the snow finally began to die down, John went through his practice routine one last time. He had gotten to the point where his nervousness had become excitement. About one hour before he was to leave, he shared a cup of tea with his roommate. They sat in the living room with the big bay window watching the snow begin to fall more heavily again. John realized suddenly that he would have to leave immediately to get there on time. As John left for his performance, his roommate said, "Break a leg tonight."

6. Ending of Prime Scenario for Experiment 2 with Literal Reference to Breaking Leg

John realized suddenly that he would have to leave immediately to get there on time. While rushing to his performance, John broke his leg.

REFERENCES

- Abelson, R. B., & Prentice, D. A. (1997). Contrast tests of interaction hypotheses. *Psychological Methods*, 2, 315–328.
- Allbritton, D. W., McKoon, G., Gerrig, R. J. (1995). Metaphor-based schemas and text representations: Making connections through conceptual metaphors. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 21, 612–625.
- Anderson, M. C., & Spellman, B. A. (1995). On the status of inhibitory mechanisms in cognition: Memory retrieval as a model case. *Psychological Review*, 102, 68–100.
- Anderson, R. C., & Ortony, A. (1975). On putting apples in bottles: A problem of polysemy. *Cognitive Psychology*, 7, 167–180.
- Austin, J. (1962) *How to do things with words*. Cambridge, MA: Harvard University Press.
- Barsalou, L. (1983). Ad hoc categories. *Memory & Cognition*, 11, 211–227.
- Bargh, J. A., Chaiken, S., Gerverder, R., & Pratto, F. (1992). The generality of the automatic attitude activation effect. *Journal of Personality and Social Psychology*, 62, 893–912.
- Bargh, J. A., Chen, M., & Burrow, L. (1996). Automaticity of social behavior: Direct effects of trait construct and stereotype activation on action. *Journal of Personality and Social Psychology*, 71, 230–244.
- Bargh, J. A., & Tota, M. E. (1988). Context-dependent automatic processing in depression: Accessibility of negative constructs with regard to the self but not others. *Journal of Personality and Social Psychology*, 54, 925–939.
- Bobrow, S. & Bell, S. (1973). On catching on to idiomatic expressions. *Memory & Cognition*, 1, 343–346.
- Chen, M., & Bargh, J. A. (1997). Nonconscious behavioral confirmation processes: The self-fulfilling consequences of automatic stereotype activation. *Journal of Experimental Social Psychology*, 33, 541–560.

- Devine, P. G. (1989). Stereotypes and prejudice: Their automatic and controlled components. *Journal of Personality and Social Psychology*, 56, 680-690.
- Dijksterhuis, A., & van Knippenberg, A. (1996). The knife that cuts both ways: Facilitated and inhibited access to traits as a result of stereotype activation. *Journal of Experimental Social Psychology*, 32, 271-288.
- Galinsky, A. D., Moskowitz, G. B., & Skurnik, I. W. (in press). Counterfactuals as self-generated primes: Priming the simulation heuristic. *Social Cognition*.
- Gernsbacher, M. A. (1990). *Language comprehension as structure building*. Hillsdale, NJ: Erlbaum.
- Gernsbacher, M. A., & Faust, M. (1990). The roll of suppression in sentence comprehension. In G. B. Simpson (Ed.), *Understanding word and sentence* (pp. 97-128). Amsterdam: North Holland.
- Gernsbacher, M. A., Keysar, B., & Robertson, R. R. (1995, November). *The role of suppression in metaphor comprehension*. Paper presented at the 36th annual meeting of the Psychonomic Society, Los Angeles, CA.
- Gibbs, R. W. (1994). *The poetics of mind: Figurative thought, language, and understanding*. Cambridge, UK: Cambridge University Press.
- Gilbert, D. T., & Hixon, J. G. (1991). The trouble of thinking: Activation and application of stereotypic beliefs. *Journal of Personality and Social Psychology*, 60, 509-517.
- Glucksberg, S., & Estes, Z. (in press). Feature accessibility in conceptual combination: Effects of context-induced relevance. *Psychonomic Bulletin and Review*.
- Glucksberg, S., & Keysar, B. (1990). Understanding metaphorical comparisons: Beyond similarity. *Psychological Review*, 97, 1-18.
- Glucksberg, S., Mcglone, M. S., & Manfredi, D. (1997). Property attribution in metaphors comprehension. *Journal of Memory and Language*, 36, 50-67.
- Glucksberg, S., Newsome, M. R., & Goldvarg, Y. (1997). Filtering out irrelevant material during metaphor comprehension. *Proceedings of the nineteenth annual conference of the Cognitive Science Society*. Mahwah, NJ: Erlbaum.
- Grice, H. P. (1975). Logic and conversation. In P. Cole & J. Morgan (Eds.), *Syntax and semantics 3: Speech acts* (pp. 41-58). New York: Academic Press.
- Hasher, L., & Zacks, R. T. (1988). Working memory, comprehension, and again: A review and a new view. In G. G. Bower (Ed.), *The psychology of language and motivation* (Vol. 22, pp. 193-225). San Diego, CA: Academic Press.
- Heider, F. (1944). Social perception and phenomenal causality. *Psychological Review*, 51, 358-374.
- Herr, P. M. (1986). Consequences of priming: Judgment and behavior. *Journal of Personality and Social Psychology*, 51, 1106-1115.
- Higgins, E. T., & Chaires, W. M. (1980). Accessibility of interrelational constructs: Implications for stimulus encoding and creativity. *Journal of Experimental Social Psychology*, 16, 348-361.
- Higgins, E. T., Rholes, W. S., & Jones, C. R. (1977). Category accessibility and impression formation. *Journal of Experimental Social Psychology*, 13, 141-154.
- Kintsch, W. (1998). *Comprehension: A paradigm of cognition*. New York: Cambridge University Press.
- James, W. (1890). *The principles of psychology* (Vol. 1). New York: Dover.
- Macrae, C. N., Bodenhausen, G. V., & Milne, A. B. (1995). The dissection of selection in person perception: Inhibitory processes in social stereotyping. *Journal of Personality and Social Psychology*, 69, 397-407.
- Martin, L. L. (1986). Set/reset: Use/disuse of concepts in impression formation. *Journal of Personality and Social Psychology*, 51, 493-504.

- McGlone, M. S., Glucksberg, S., & Cacciari, C. (1994). Semantic productivity and idiom comprehension. *Discourse Processes*, 17, 167–190.
- Moskowitz, G. B., & Roman, R. J. (1992). Spontaneous trait inferences as self-generated primes: Implications for conscious social judgment. *Journal of Personality and Social Psychology*, 62, 728–738.
- Moskowitz, G. B., Skurnik, I. & Galinsky, A. D. (1999). The history of dual process notions and the future of preconscious control. In S. Chaiken & Y. Trope (Eds.), *Dual process theories in social psychology* (pp. 12–36). New York: Guilford.
- Moskowitz, G. B., & Skurnik, I. (1999). The differing nature of contrast effects following exemplar versus trait primes. Standard of comparison versus motivated correction processes. *Journal of Personality and Social Psychology*, 76, 911–927.
- Ortony, A. (1975). Why metaphors are necessary and not just nice. *Educational Theory*, 26, 395–298.
- Rosch, E. (1978). Principles of categorization. In E. Rosch & B. B. Lloyd (Eds.), *Cognition and categorization* (pp. 27–48). Hillsdale, NJ: Erlbaum.
- Rothbart, M., Sriram, N., & Davis-Stitt, C. (1996). The retrieval of typical and atypical category members. *Journal of Experimental Social Psychology*, 32, 309–336.
- Sedikides, C., & Skowronski, J. J. (1991). The law of cognitive structure activation. *Psychological Inquiry*, 2, 169–184.
- Simpson, G. B. (1984). Lexical ambiguity and its role in models of word recognition. *Psychological Bulletin*, 96, 316–340.
- Simpson, G. B. (1994). Context and the processing of ambiguous words. In M. A. Gernsbacher (Ed.), *Handbook of psycholinguistics* (pp. 359–374). San Diego, CA: Academic Press.
- Simpson, G. B., & Kang, H. (1994). Inhibitory processes in recognition of homograph meanings. In D. Dagenbach & T. H. Carr (Eds.), *Inhibitory processes in attention, memory and language* (pp. 359–378). San Diego, CA: Academic Press.
- Simpson, G. B., & Krueger, M. A. (1991). Selective access of homograph meanings in sentence context. *Journal of Memory and Language*, 30, 627–643.
- Srull, T. R., & Wyer, R. S., Jr. (1980). Category accessibility and social perception: Some implications for the study of person memory and interpersonal judgment. *Journal of Personality and Social Psychology*, 38, 841–856.
- Stapel, D. A., Koomen, W., & van der Pligt, J. (1996). The referents of trait inferences: The impact of trait concepts versus actor-trait links on subsequent judgments. *Journal of Personality & Social Psychology*, 70, 437–450.
- Stapel, D. A., Koomen, W., & van der Pligt, J. (1997). Categories of category accessibility: The impact of trait concept versus exemplar priming on person judgments. *Journal of Experimental Social Psychology*, 33, 47–76.
- Strack, F., Schwarz, N., Bless, H., Kübler, A., & Wänke, M. (1993). Awareness of the influence as a determinant of assimilation versus contrast. *European Journal of Social Psychology*, 23, 53–62.
- Uleman, J. S., Newman, L. S., & Moskowitz, G. B. (1996). People as flexible interpreters: Evidence and issues from spontaneous trait inference. In M. Zanna (Ed.), *Advances in experimental social psychology* (Vol. 28, pp. 211–280). San Diego, CA: Academic Press.
- Wegener, D. T., & Petty, R. E. (1995). Flexible correction processes in social judgment: The role of naive theories in corrections of perceived bias. *Journal of Personality and Social Psychology*, 68, 36–51.

Reproduced with permission of the copyright owner. Further reproduction prohibited without permission.