

2

[illegible]

1/ Nov 92

ANNUAL REPORT - 3/1/91 - 2/28/92

Discourse Models, Pronoun Resolution, and the Implicit Causality of Verbs

Gail McKoon; Steven B. Greene & Roger Ratcliff

PE 61102F
PR 2313
TA A4
GR - AFOSR-90-0246

Northwestern University
Psychology Department
Evanston IL 60208

AFOSR-TR-82-040

Air Force Office of Scientific Research/NL
Building 410
Bolling AFB DC 20332-6448

Approved for public release;
distribution unlimited

12. ABSTRACT (Maximum 200 words)

Some interpersonal verbs, such as *admire* and *amaze*, describe an action or property of one person (the reactor) that is necessarily a response to an action or property of another (the initiator). We hypothesized that these verbs make the initiator relatively more accessible in a comprehender's discourse model, and that this change in relative accessibility would aid identification of the referent of a pronoun in a subsequent *because* clause. We predicted that, as a result, subjects would be faster to recognize a character's name after a *because* clause that uses a pronoun that refers to that character than after one that refers to the other character. Four experiments confirmed this prediction. Three further experiments demonstrated the importance of the verb's causal structure and of the presence of the connective *because* to this result.

11. NUMBER OF PAGES: 10

1000000.COM

UNCLASSIFIED

UNCLASSIFIED

UNCLASSIFIED

UNLIMITED

Discourse Models, Pronoun Resolution, and the Implicit Causality of Verbs

Gail McKoon

Northwestern University

Steven B. Greene

Princeton University

Roger Ratcliff

Northwestern University

DTIC QUALITY INSPECTED 5

Running head: Implicit Causality of Verbs

Address correspondence to:

Gail McKoon
Psychology Department
Northwestern University
Evanston, IL 60208

Accession For	
NTIS	CRA&I ☒
DTIC	TAB ☐
Unannounced	☐
Justification	
By	
Distribution /	
Availability Codes	
Dist	Avail and/or Special
A-1	

93-01472



93 1 26 035

17 NOV 1992

Abstract

Some interpersonal verbs, such as *admire* and *amaze*, describe an action or property of one person (the reactor) that is necessarily a response to an action or property of another (the initiator). We hypothesized that these verbs make the initiator relatively more accessible in a comprehender's discourse model, and that this change in relative accessibility would aid identification of the referent of a pronoun in a subsequent *because* clause. We predicted that, as a result, subjects would be faster to recognize a character's name after a *because* clause that uses a pronoun that refers to that character than after one that refers to the other character. Four experiments confirmed this prediction. Three further experiments demonstrated the importance of the verb's causal structure and of the presence of the connective *because* to this result.

The use of psychological methods to study linguistic phenomena offers the possibility of simultaneous progress on issues in both fields. At least as far back as early empirical investigations of the derivational theory of complexity (e.g., Fodor & Bever, 1965; Fodor, Garrett, & Bever, 1968; Miller, 1962), psychologists have sought empirical evidence for hypotheses put forth by their colleagues in linguistics. The finding of such evidence both supports the linguistic hypotheses and allows the construction of models of underlying psychological processes that presumably rely on linguistic regularities.

In what follows, we describe the use of psychological methods to study the processes of pronoun resolution during comprehension of linguistic stimuli of special interest. These stimuli are of special interest because they employ verbs from a class exhibiting "implicit causality" (Garvey & Caramazza, 1974). We specify the nature of this implicit causality in greater detail below; for now, some illustrations will make this property clear. Consider the sentence frame "Mathilda amazed Jonathan because...." When asked to complete a sentence frame of this form, subjects show great regularity in choosing to say something about Mathilda rather than about Jonathan. Note that either type of continuation is possible, e.g., "because she displayed such refined talent" or "because he had never seen a fire-eater before." Garvey and Caramazza identified this type of implicit causality as NP₁ causality because the bias is to continue the sentence by saying something about the surface subject. Some verbs exhibit NP₂ causality instead, such as in "Felix admired Alexandra because....," which most subjects will complete by describing a property of Alexandra's ("because she aced the accounting exam") rather than a property of Felix's ("because he was always in desperate need of a role model"). A number of verbs exhibit NP₁ causality; a number of others exhibit NP₂ causality. We discuss below the characteristics of these two groups of verbs.

Psychologists studying language have long been interested in how information conveyed by the main verb of a sentence contributes to the sentence's grammatical structure (e.g., Healy & Miller, 1971). More recently, their attention has focused on the particular issue of the implicit

causality of verbs, which has been studied using a variety of tasks (Au, 1986; Brown & Fish, 1983; Caramazza, Grober, Garvey, & Yates, 1977; Ehrlich, 1980; Hoffman & Tchir, 1990; Hudson, Tanenhaus, & Dell, 1986). However, there has to date been no systematic, empirical demonstration that implicit causality is understood except under conditions in which subjects have been asked to engage in some explicit strategy; for example, they may be asked to generate a continuation for the sentence or to speak aloud the name of a character referred to by a pronoun. Whether implicit causality is understood in the absence of such specific strategies is still an open question. Ideally, we would like an empirical demonstration that implicit causality has an effect on comprehension, plus some method for measuring that effect. One place we might observe an effect of implicit causality is in the processes that identify an argument of a verb as the referent of a subsequent pronoun. Fortunately, there is a widely accepted technique for studying these processes—comparing the accessibility of referents and nonreferents after pronouns are read (Chang, 1980; Corbett & Chang, 1983; Dell, McKoon, & Ratcliff, 1983; Gernsbacher, 1989; MacDonald & MacWhinney, 1990; McKoon & Ratcliff, 1980, 1984).

A demonstration of effects of a verb's implicit causality on pronoun resolution would be especially interesting in light of the difficulty of finding evidence of pronoun resolution in other contexts. Recently, Greene, McKoon, and Ratcliff (1992) proposed a framework in which to study pronoun processing. According to the Greene et al. framework, comprehenders construct a discourse model that represents the entities and events evoked by a discourse and the relationships among them (see Grosz, 1981; Grosz, Joshi, & Weinstein, 1983; Grosz & Sidner, 1986; McKoon, Ratcliff, Ward, & Sproat, 1992; McKoon, Ward, Ratcliff, & Sproat, 1992; Sidner, 1983a, 1983b; Ward, Sproat, & McKoon, 1991; Webber, 1983). Each entity in the discourse model has some degree of accessibility relative to all other entities. The initial degree of accessibility of an entity is determined by the syntactic, semantic, and pragmatic means by which it is introduced, and its accessibility changes as comprehension of various syntactic and semantic structures alters the relationships represented in the model. The accessibility of an

entity in a discourse model is therefore determined not only by the manner in which it is introduced into the discourse, but also by subsequent references to it.

In this framework, the job a pronoun performs is seen not as a trigger that initiates a serial search for an antecedent (see Matthews & Chodorow, 1988), but as a cue to identify the discourse entity that best matches the semantic and grammatical features of the pronoun. Specifically, the identification of a referent for a pronoun is first attempted by a fast, automatic process that depends on the accessibility of the intended referent in the discourse model. This process matches the features of the pronoun in parallel against those of all entities in the discourse model. If one entity matches sufficiently well and better than all other entities, it is identified as the most likely referent of the pronoun. On the other hand, if either no referent matches sufficiently or more than one referent matches equally well, the comprehender may optionally engage in further, strategic, processing to identify the referent. A series of experiments by Greene et al. in which subjects read short (three-sentence) texts describing two equally salient characters found evidence of successful pronoun resolution only when subjects had extrinsic motivation to keep track of the characters and generous time in which to do so. In the absence of these factors, no evidence of pronoun resolution was found. The pronoun-as-cue framework explains this result: Because the two entities were equally salient, neither matched the pronoun sufficiently better than the other to be uniquely identified as its likely referent. Based on this evidence, Greene et al. argued that the processes responsible for pronoun resolution in previous psychological experiments (e.g., Chang, 1980; Corbett & Chang, 1983; Gernsbacher, 1989) have been optional, strategic processes and not a mandatory component of comprehension.

In contrast to typical experimental materials that describe two characters who are equally in the focus of attention, natural discourse commonly uses a pronoun to refer to a discourse entity that is already highly salient, relative to other entities (Brennan, 1989; Chafe, 1974; Ehrlich, 1980; Fletcher, 1984; Greene et al., 1992; see also Givon, 1976). The occurrence of a

pronoun usually indicates to the comprehender that the discourse is still centered on the previously salient entity or entities (Greene et al., 1992; Grosz et al., 1983). Numerous syntactic, semantic, and pragmatic devices can be used to establish one discourse entity as the current focus of attention and, therefore, as likely to be referred to subsequently (Gernsbacher, 1990; Gernsbacher & Shroyer, 1989; Grosz, 1981; McKoon, Ratcliff, Ward, & Sproat, 1992; McKoon, Ward, Ratcliff, & Sproat, 1992; Sidner, 1983a; Ward et al., 1991). An utterance containing a verb exhibiting implicit causality may have the effect of establishing the verb's more prominent argument as the current focus of attention (Hudson et al., 1986). In the terms of the pronoun-as-cue framework, these verbs may alter the relative accessibilities of their arguments in a discourse model. That change in accessibility may be sufficient to ensure that the fast, automatic process of pronoun resolution can provide one of them as the likely referent of a subsequent pronoun. If that is the case, then we may be able to find evidence of successful pronoun resolution even when the experimental procedures employed do not encourage subjects to engage in strategic processing.

Before turning to the empirical evidence, let us examine in greater detail why some verbs exhibit the implicit causality that we hypothesize to privilege one possible referent over the other in a discourse model framework. Garvey and Caramazza (1974) coined the term "implicit causality" to describe a property of transitive verbs that relate two nouns referencing human or animate beings in such a way that "[o]ne or the other of the noun phrases is implicated as the assumed locus of the underlying cause of the action or attitude" (p. 460). Garvey and Caramazza argued that implicit causality is part of the semantics of the verb root: Some verbs, such as *confess*, *telephone*, and *approach*, assign the cause of the event to the subject noun phrase (NP₁), while others, such as *fear*, *praise*, and *admire*, assign the cause to the object noun phrase (NP₂). By examining subjects' completion of sentence frames such as, "The prisoner confessed to the guard because he . . .," these researchers established that, when asked to do so, English speakers reliably attribute causality to NP₁ for some verbs and to NP₂ for other verbs.

A subsequent experiment (Caramazza et al., 1977) revealed that subjects were faster to name the antecedent for a pronoun after reading a sentence containing a verb exhibiting implicit causality if that pronoun was consistent with the causality than if it was not. For example, when asked to identify the referent for *he*, subjects responded "Jimmy" faster after reading *Jimmy confessed to Mary because he wanted forgiveness* than they responded "Michael" after reading *Cathy confessed to Michael because he offered forgiveness*.

Garvey and Caramazza (1974) identified the "locus of the underlying cause" as the relevant factor in determining a verb's implicit causality, but they stopped short of a full explanation of why that factor is critical and how one determines this locus. Following Au (1986; also Osgood, 1970), we propose to analyze interpersonal verbs in terms of which of their arguments initiates a state of affairs and which one reacts to it. The implicit causality of a verb is towards the argument that initiates an action or evokes a response. As noted earlier, the subject of *confess* initiates the action: We confess for things we ourselves have done. In contrast, the subject of *thank* is reacting to a state of affairs brought about by the object: We thank others for things they have done. In one case, the grammatical subject is the initiator, and the object is the reactor; in the other, the object is the initiator, and the subject is the reactor. Note that the reactor may very well carry out some action, as in *thank*, as well as in *correct* and *congratulate*; the key is that the action is necessarily in response to an initiating state or action of someone else. Often the reactor's action is a speech act, but it need not be, as in *help*.

Levin's (1992) recent discussion of English verb classes supports the initiating/reacting distinction. Levin, summarizing earlier work in linguistics, classifies verbs of psychological states ("psych-verbs"), such as *amaze* and *admire*, into two categories, depending on whether the experiencer of some emotional reaction is the surface subject or object. She also describes another category, "judgment verbs," such as *congratulate*, *reproach*, and *scold*, which are like the *admire* psych-verbs in that those verbs "relate to a particular feeling which someone may have in reaction to something, [and] these verbs relate to a judgment or opinion which someone

may have in reaction to something" (p. 175). Thus, both the *admire* verbs and the judgment verbs indicate that the surface subject is experiencing some reaction at the initiation of the surface object. Levin's analysis of judgment verbs is reminiscent of Fillmore's (1971) analysis of the same verbs as presupposing responsibility on the part of the argument filling the role he labeled "defendant," generally the surface object.

The initiating/reacting distinction intuitively matches our understanding of implicit causality. Subjects' completions of *because* clauses reveal what aspect of the verb's meaning subjects believe requires a causal explanation. The initiating of a state of affairs typically demands an explanation; the reaction is explained by the state of affairs itself. Thus, *because* clauses should typically explain the behavior of the initiator, not the reactor.

To summarize, verbs that exhibit implicit causality are those whose arguments fill the roles of initiator and reactor. Some property or action of the initiator causes a response by the reactor; this response may simply be an emotion (*admire*) or a perception (*notice*), or it may include an action (*thank*). A *because* clause will naturally then explain what property or action of the initiator provoked the response by the reactor. However, as Garvey and Caramazza (1974) first noted, it is, of course, possible for *because* clauses to offer an explanation in terms of a property or action of the reactor, as in *Cathy confessed to Michael because he offered forgiveness*. In such an instance, in which the *because* clause is inconsistent with the implicit causality of the verb, the analysis requires an additional step. A property or action of the initiator still causes a response by the reactor, but the nature of the explanation offered by the *because* clause is different. In this case, the *because* clause explains what property or action of the reactor made the initiator's property effective or the initiator's action possible.

Although our analysis of implicit causality is compatible with current linguistic discussions of the argument-taking properties of verbs, it differs somewhat from that found in previous psychological work (e.g., Brown & Fish, 1983). Researchers since Garvey and Caramazza's original work have sometimes replaced their atheoretical NP_i/NP_r classification

scheme with one that distinguishes between "state verbs," which describe a situation in which one person (the stimulus) induces a psychological state in another (the experiencer), and action verbs, which describe a situation in which one person (the agent) instigates an action directed at another (the patient) (Brown and Fish, 1983). According to Brown and Fish's analysis, state verbs will exhibit implicit causality for NP₁ or NP₂, depending upon which noun phrase refers to the stimulus. Action verbs, in contrast, should always exhibit implicit causality for NP₁, the agent, according to this analysis. However, Au (1986) found that although some action verbs, such as *cheat* and *flatter*, exhibit implicit agent causality, others, such as *correct* and *praise*, exhibit implicit patient causality. Au instead resurrected an earlier analysis of causal attribution, that of Osgood (1970), to explain the implicit causality of action verbs, while retaining the Brown and Fish analysis of state verbs.

Our conclusion is that the state/action distinction is superfluous to understanding implicit causality. Implicit causality has been found to be a property of some, but not all, verbs in both categories. Therefore, classifying a verb as belonging to either category tells us little about whether that verb will exhibit implicit causality, and, further, classifying a verb as an action verb tells us nothing about which way the causality will go. No matter whether we categorize a verb as action or state, we still have to analyze its semantics further to predict its implicit causality. We instead analyze both state and action verbs solely in terms of the initiating and reacting roles of their arguments to predict implicit causality.

Experiments 1, 2, 3, and 4

Table 1 shows examples of the texts that were used in the experiments. Consider the first example in Table 1; in the third sentence, *infuriate* is a verb for which the subject—in this case, James—is the initiator. The subject does something or has some property that brings about a reaction by the object; in this case, the reaction is an emotion. The example shows two possible

continuations of the third sentence: In the first, the *because* clause is consistent with the implicit causality of *infuriate*; in the other, it is inconsistent. Given our analysis of verbs exhibiting implicit causality and the pronoun-as-cue processing hypothesis, we can suggest how the two alternative continuations of the final sentence might be understood during comprehension. As a verb exhibiting implicit causality, *infuriate* makes the initiator, James, relatively more accessible than other entities in the discourse model of the text. In the first continuation, *he leaked important information to the press*, the pronoun is intended to refer to James. When it is matched as a cue against the entities in the discourse model, the most accessible entity, James, is returned as the most likely referent. The gender of the pronoun is consistent with James as the referent, and perhaps more importantly, the information in the continuation is consistent with the implicit causality structure of the verb; it explains what state of affairs James created. The several factors of increased accessibility in the discourse model, gender agreement, and appropriateness of the continuation for the verb's causality all conspire toward identification of James as the referent for the pronoun.

Insert Table 1 about here

In contrast, consider the second continuation, *she had to write all the speeches*. The most accessible referent is still the initiator, James, but now the gender of the pronoun does not match. Moreover, the content of the continuation is inconsistent with the verb's implicit causality. The predicate explains what Debbie had to do in response to the state of affairs created by James, not what James himself did. Because of these mismatches, the initiator should be discarded as a potential referent. The remaining two possibilities are that pronoun resolution may fail, leaving the pronoun reference unresolved, or that the other, intended, referent—Debbie—may be selected.

The situation is similar for verbs for which the object is the initiator, like *value*, in the second example in Table 1. The object of *value* does something or has some property that brings

about a reaction by the subject. Thus, *value* makes its object relatively more accessible in a discourse model. In the first continuation, *he always knew how to negotiate*, which is consistent with the implicit causality of *value*, the pronoun is intended to refer to Sam, and the continuation explains what property of Sam's prompted Diane's reaction. So, when the pronoun is matched against the discourse model, Sam is returned as the most likely referent, and the matching gender and consistent continuation confirm this selection.

Once again, in the other continuation, *she never knew how to negotiate*, the pronoun mismatches the most accessible entity on gender, and the information in the continuation is inconsistent with the causality implicit in the verb. The continuation explains what property of Diane's allowed her to appreciate the property of Sam's, and only indirectly what property Sam possessed. As with the inconsistent continuation of the subject-initiating verb *infuriate*, pronoun resolution may fail, or the only other referent, the reactor, may be selected.

All of the experiments described here compared subjects' reaction times to recognize a character's name as having appeared in the current text when the subjects were tested after the two types of continuations: those in which a pronoun referred to the tested character and those in which a pronoun referred to the other character. The test always occurred at the end of the third sentence of three-sentence texts like those in Table 1. Following the reasoning just outlined, for the character that was the referent of the pronoun in the consistent continuation (e.g., James in the first example in Table 1), we anticipated that responses to that character's name would be facilitated when it was tested after the consistent continuation relative to the inconsistent continuation; that is, responses would be facilitated for the name when that character was the referent versus when it was not. We refer to this as a matching effect: Responses to a character's name are facilitated when it matches the referent of the pronoun versus when it does not.

However, for the character intended as the referent in the *inconsistent* continuation two outcomes are possible. In this case, the processes of pronoun resolution may leave the reference

unresolved, resulting in no matching effect, but perhaps overall facilitation for the initiator role because of its initial greater accessibility. Or, if the pronoun resolution process does not fail but instead selects the other character, the reactor, as the referent for the pronoun, we would again expect facilitation for the character referred to by the pronoun, in this case, the reactor. We would therefore expect a matching effect such that responses are facilitated when the character's name presented for recognition matches the referent of the pronoun in the continuation.

Experiments 1 and 2 examined subject-initiating verbs, like *infuriate*, and Experiments 3 and 4 examined object-initiating verbs, like *value*. These experiments were designed to examine pronoun resolution under conditions in which subjects read at approximately normal rates without adopting any special strategies. The materials were presented at a rate of about 250 ms/word, a rate that other research (e.g., Dell et al., 1983, Greene et al., 1992, Experiments 8 and 9; Just & Carpenter, 1980; Rayner, 1978) has shown to be reasonable for college students. Comprehension questions following the texts asked about a variety of information from the texts; they did not ask about specific kinds of information, such as which character carried out particular actions, so as not to induce subjects to adopt strategies specific to pronoun resolution (or any other task beyond that required by the experimental procedure directly). Finally, three times as many filler items as critical items were included in the experiments in order to reduce the predictability of the type of item to be tested and the test locations.

Method

Materials. Twenty subject-initiating verbs and 20 object-initiating verbs were chosen from those used in previous research (Au, 1986; Brown & Fish, 1983). Because we selected only verbs that were subject- or object-initiating according to our analysis of implicit causality, we excluded some verbs, such as *telephone* and *hit*, that had been included in previous research. The subject-initiating verbs we selected were *aggravate*, *amaze*, *amuse*, *annoy*, *apologize*, *bore*,

charm, cheat, confess, deceive, disappoint, exasperate, fascinate, frighten, humiliate, infuriate, inspire, intimidate, scare, and surprise. The object-initiating verbs were *assist, blame, comfort, congratulate, correct, detest, dread, envy, hate, help, jeer, notice, pacify, praise, reproach, scold, stare, thank, trust, and value.* The implicit causality of these verbs can be demonstrated by asking subjects to generate continuations of sentence fragments that present the verbs in the frame "proper noun, verb (tense), proper noun, *because*" (e.g., *James infuriated Debbie because____*). Continuation data were collected for some of the 40 verbs used in our experiments by Au (1986), and we collected continuation data for the others. Overall, the mean percentage of subjects continuing a sentence fragment with a pronoun referring to the referent consistent with the causality of the verb was 89 for the subject-initiating verbs and 92 for the object-initiating verbs.

Each verb was used in the third sentence of a three-sentence text. The first sentence of each text introduced two characters, one male and the other female, and the third sentence mentioned these characters again by name. The second sentence referred to both of them by anaphora (usually, *they*). For half of the texts, the *first-mentioned* character in both the first and third sentences was male and for the other half, female. The critical verb was used in the first clause of the third sentence. The two clauses of the third sentence were always joined by *because*. There were two versions of the second clause of the third sentence: One version began with a pronoun matching the gender of the first character in the first clause and continued with information that made sense for that character in a causal role; the second version began with a pronoun matching the gender of the other character, and continued with information that made sense for that character. An example of a text for a verb with each kind of implicit causality is shown in Table 1. The average length of the first and second sentences combined was 19.8 words, and the average length of the third sentence was 10.9 words. The average number of words between the first character's name in the first clause of the third sentence and the pronoun in the second clause was 3.2; the average number of words between the second character's name

and the pronoun was 1 (*because*), and the average number of words between the pronoun and the end of the sentence was 5.7. There were two test words for each text, the two character names. There were also two test statements for each text, one true and one false. These tested a variety of kinds of information from the texts.

There were 60 filler texts used to provide different kinds of test words from the experimental texts. These texts were all three sentences long and averaged 33 words in length. Each text had one test word. Thirty-five of these test words had not appeared in any text (17 of these were proper names), and 25 had appeared in their text. Nineteen were tested in the first two sentences, and the remainder were tested in the third sentence. Each filler text had associated with it one true and one false test statement; as with the experimental texts, these were written to test a variety of kinds of information from the texts.

Procedure. All of the texts and test items were presented on a CRT screen, and responses were collected on the CRT keyboard. Each subject participated in one 50-minute session.

Each experiment began with 30 lexical decision test items. These items were included to give subjects practice with the response keys on the CRT keyboard. After this practice, there were 20 filler texts, and then the remainder of the texts—20 experimental (20 subject-initiating texts in Experiments 1 and 2, and 20 object-initiating texts in Experiments 3 and 4) and 40 fillers—were presented in random order.

Each text began with the instruction to press the space bar on the keyboard to initiate the text. When the space bar was pressed, the text was presented, one word at a time. Each word was displayed in the same location on the CRT screen, and each was displayed for 170 ms plus 17 ms multiplied by the number of letters in the word. There was no pause between words. The last word of a sentence was displayed for an extra 200 ms unless it was immediately followed by a test word. When a test word was presented, it appeared in the same location as the text words; its letters were all in upper case (unlike the words of the text) and two asterisks were displayed

immediately to its left and to its right. The test word remained on the screen until a response key was pressed (y for “yes” the word had appeared in the text, and n for “no” the word had not appeared in the text). After the response and a pause of 170 ms, the text continued or the “press space bar” message for the next text was presented, except, in Experiments 2 and 4, if the response was slower than 1100 ms, the message *TOO SLOW!* was displayed first for 500 ms. Each text was followed by a true/false test statement, and incorrect responses to this test statement were followed by an error message, the word *ERROR* presented for 1500 ms. Each text had a true and a false test statement; which one of these was presented was chosen randomly. For the test words, subjects were instructed to respond as quickly and accurately as possible. For the true/false test statements, they were told to aim for 100% accuracy. Experiments 2 and 4 used the “too slow” message to speed subjects’ responses because we wanted to be sure that their decisions about the test words could not be based on slow, strategic processes that began at the time of the presentation of the test word.

Design and Subjects. For all four experiments, there were two variables for the 20 experimental texts: the pronoun in the second clause of the third sentence matched in gender either the first character in the first clause or the second character, and the test word was either the name of the first character or the second. Note that the consistent pronoun refers to the first character name for the subject-initiating verbs and the second character name for the object-initiating verbs. For the experimental texts, the test word was always presented after the final word of the text. The four conditions formed by crossing the two variables were combined in a Latin square design with four sets of texts (5 per set) and four groups of subjects (5 in each group except for Experiment 2, in which there were 7 in each group). The subjects participated in the experiments for credit in an introductory psychology course.

Results and Discussion

Means were calculated for each subject and each item in each condition, and means of these means are shown in Table 2. All response times longer than 2000 ms were eliminated from the means and analyses. For Experiments 1 and 3, this was about 4% of the data, and for Experiments 2 and 4, this was almost none of the data. Response times for filler test words and true/false test statements are shown in Table 6 for all the experiments.

Insert Table 2 about here

Examination of the data in Table 2 shows that which pronoun was used in the text had a strong effect on response times to the test words. Consider, for example, Experiment 1, in which the consistent pronoun referred to the first character mentioned in the final sentence. When the consistent pronoun was used in the sentence, responses to the first character name as a test word were faster than when the inconsistent pronoun was used. In other words, responses to the test word were faster when the test word matched the referent of the pronoun than when it did not. A similar matching effect was obtained when the second character name was presented as a test word: When it matched the referent of the pronoun, responses were faster than when it did not match. We interpret the matching effect as showing that the subjects in these experiments understood which of the two characters in a text was the intended referent of the pronoun, in contrast to previous experiments in which they did not (Greene et al., 1992).

However, one caveat is in order. It should be clear that we have no measure of a neutral baseline for response times to our recognition tests of the characters' names following the texts. In the experiments in Greene et al., we used sentences like *Mary accidentally scratched John with a knife and then she dropped it on the counter*. We measured the response time to a character's name both before and after the pronoun in the second clause of its sentence, so that

we could examine the relative facilitation given by the pronoun to its referent versus a non-referent. Whether any obtained facilitation was due to true facilitation for the referent or inhibition for the non-referent is impossible to determine. Similarly, in the experiments reported here, we compared whether the response time to a character's name at the ends of the sentences changed as a function of whether the name matched the referent of the pronoun in the sentence, but whether that change was facilitation for a referent or inhibition for a nonreferent is impossible to say. Because we were concerned only with relative effects, this is not a serious problem. Our claim is only that the matching effect represents a relative change in the accessibilities of the referent versus the nonreferent.

The matching effect held for both subject-initiating verbs and object-initiating verbs, and for subjects who were pressed to respond quickly (by the "too slow" message) and those who were not, with one exception. For the subject-initiating verbs tested with the "too slow" message, the test word corresponding to the referent of the consistent pronoun did not show a matching effect. In this one case, response times did not appear to slow significantly when the test word did not match the referent of the pronoun, and this result suggests that pronoun resolution may be somewhat less robust with subject-initiating verbs than with object-initiating verbs.

The pattern of data for each experiment represents an interaction between which character name was tested and which pronoun was used in the sentence. The significance of the interactions was demonstrated by F_1 (subjects) and F_2 (items) values 5.0 or larger in all four experiments, by analyses of variance that treated subjects as the random variable and analyses that treated items as the random variable. In Experiment 2, there was one additional significant effect: The first test word had faster response times than the second test word, $F_1(1,27)=14.2$ and $F_2(1,19)=5.8$. No other effects approached significance in both subjects and items analyses. Error rate differences were also tested by analyses of variance, and all F values were not significant ($p > .05$, F s less than 3.1), except in Experiment 3, where there were significantly

more errors on the second test word than the first, $F_1(1,19)=5.9$ and $F_2(1,19)=4.1$. Standard errors of the response time means are shown in Table 6 (for all experiments).

Experiments 5 and 6

Experiments 1 through 4 demonstrated a matching effect such that responses to a test of a character's name were facilitated if the character matched the referent of the preceding pronoun. We have hypothesized that this happened because the structure of verbs exhibiting implicit causality privileges the initiator role over the reactor role as a potential pronominal referent. If the gender of the subsequent pronoun and the information in the continuation following the pronoun are consistent with the implicit causality of the verb, the character in the initiator role is taken to be the pronoun's referent, as demonstrated by the matching effect observed for the consistent continuations in Experiments 1 through 4. If, however, the gender of the pronoun and the information in the predicate are inconsistent with the potential referent privileged by the verb's implicit causality, this mismatch causes the other character, the reactor, to be selected as the referent of the pronoun, as demonstrated by the matching effect for the inconsistent continuations. For both consistent and inconsistent continuations, the result is the same—faster recognition responses to a character's name if that character matches the referent of the pronoun in the continuation.

Our account of the matching effects found in Experiments 1 through 4 emphasizes the importance of consistency between the verb's causal structure and the explanation of the verb's action given in the *because* clause. The relationship between the two is made explicit by the word *because*. This connective may serve to bring to the fore the information about implicit causality inherent in the verb's lexical structure. Experiments 5 and 6 examine whether the presence of this connective is necessary to create the effect observed in Experiments 1 through 4.

Method

Experiment 5 examined subject-initiating verbs, and Experiment 6 examined object-initiating verbs. The 20 texts for the subject-initiating verbs and the 20 texts for the object-initiating verbs were each modified so that the final, two-clause, sentence became two sentences with *because* deleted. This was the only change made to the materials. For example, the final sentences for the first text in Table 1 were changed to: *James infuriated Debbie. He leaked important information to the press*, and *James infuriated Debbie. She had to write all the speeches*. As these examples suggest, it is still possible, or even likely, that comprehenders will interpret the information in the second sentence as a reason for the action in the first sentence. However, the relation is not made explicit in the text; instead comprehenders must make what Clark (1977) refers to as a bridging inference. We hypothesized that less causally explicit materials might adversely affect pronoun resolution, causing the matching effect to be reduced or to disappear altogether. Of course, changing the two clauses of the original version of the sentence into two separate sentences may also modify discourse relations in other ways, but we lack a sufficiently thorough understanding of discourse representation to predict such changes with any precision. Hence, our interpretations of results from this experiment will of necessity be tentative.

In displaying the two final sentences, the words were presented as in the previous experiments, and there was an additional 200 ms pause after the final word of the first of the two sentences. In all other respects, the experimental procedures and materials were the same as in the previous experiments. (There were no "too slow" messages.) The test words for the experimental texts were always presented at the end of the final sentence of their text. There were the same two variables as in the previous experiments: The final sentence used either the consistent or the inconsistent pronoun, and the test word was either the first character's name or the second character's name. These four conditions were combined in a Latin square design,

with 28 subjects in each experiment.

We also collected continuation data on these new materials. We wondered whether the same preference to refer to either the surface subject or surface object shown in continuations with *because* sentences would also appear without the *because* connective. For the continuation study, we modified the two final sentences of each text so that they used two names of the same gender, and presented them in the frame: "proper name, verb (tense), proper name, pronoun" (e.g. *James infuriated Sam. He*____). Subjects were asked to continue the second sentence, and their continuations were scored according to whether the content indicated that the pronoun had been interpreted as referring to the first character or the second. The texts were divided into two sets, each with half subject-initiating verbs and half object-initiating verbs randomly ordered, and 42 subjects gave continuations for each set. For the subject-initiating verbs, the probability of a continuation indicating that the pronoun had been interpreted according to the causality of the verb was high, .88, as it had been with the connective *because*. But for the object-initiating verbs, the preference was no longer evident; the probability of a continuation indicating interpretation of the pronoun according to the causality of the verb was only .39. These proportions most likely indicate a preference for a subsequent sentence to refer to the surface subject of a preceding sentence.

Results and Discussion

The data were analyzed as for the previous experiments (with responses slower than 2000 ms, less than 2%, eliminated), and means are shown in Table 3.

Insert Table 3 about here

The only difference between these two experiments, 5 and 6, and Experiments 1 and 3 was that the connective *because* was deleted, turning the two-clause final sentences of Experiments 1 and 3 into two separate sentences in Experiments 5 and 6. This difference eliminated the matching effect completely; in Experiments 5 and 6, response time for a test word was not affected by whether or not it matched the intended referent of the pronoun that preceded it. In fact, the only effect in response times was that, for the object-initiating verbs, responses to the first character name (the name that the pronoun would not be expected to match) were slower than responses to the second character name. This effect was significant, $F_1(1,27)=4.5$ and $F_2(1,19)=5.5$. All other F s, for both experiments, were less than 1.0. There were no significant effects on error rates, F s less than 1.5.

Clearly, the presence of the connective *because* contributes to successful pronoun resolution following verbs exhibiting implicit causality. This finding suggests that the lexical structure of the verb and the information contained in the sentence continuations are not sufficient either alone or in combination to bring about successful pronoun resolution. Instead, the relation between the continuation and the causal structure in the verb's lexical representation must be made explicit as was done in Experiments 1 through 4 with the connective *because*.

Experiment 7

Experiments 1 through 4 found evidence of facilitation for test words that match the referent of the preceding pronoun in a *because* clause following verbs that exhibit implicit causality. Experiments 5 and 6 suggested that the *because* connective is critical to this matching effect. This suggests a further possibility to be examined: Perhaps the presence of *because* is not only necessary, but, in fact, is sufficient to create the effect. The results obtained in Experiments 1 through 4 were obtained using materials with *because* connectives; earlier

failures to find similar evidence of pronoun resolution used materials with no *because* clauses (Greene et al., 1992). This final experiment examines whether adding *because* clauses to those earlier materials might allow us to find evidence of pronoun resolution.

Method

Materials. The 32 experimental texts were modified from texts previously used by Greene et al. (1992). An example text is shown in Table 4. Each text was made up of three sentences, with the first sentence introducing two characters of different genders and the second sentence referring to both of them anaphorically. There were two versions of the third sentence, each made up of two clauses connected by *because*. The first clause was the same in both versions and mentioned both characters by name, in the same order as in the first sentence. The first name was the subject of the verb in this clause; the second name was usually a direct or indirect object. None of the verbs used in these clauses fit our analysis of verbs that exhibit implicit causality. One of the second clauses of the final sentence referred to the first character with a pronoun and continued with information consistent with that character in a causal role. The other second clause referred to the second character with a pronoun and continued with information consistent with that character. The mean number of words in the first two sentences was 18.2; the mean number of words in the third sentence was 14.0. The mean number of words between the first character's name in the third sentence and the pronoun was 7.1; between the second character's name and the pronoun, 2.2. The mean number of words between the pronoun and the end of the sentence was 4.9. There were two test words for each text, the two character names. There was one true/false test statement for each text; half were true and half false. The same filler texts were used as in the previous experiments.

Insert Table 4 about here

We collected continuation data for the final sentences of these texts in the same way as for the texts used in Experiments 1 through 4. The first clause of each final sentence plus the word *because* was presented as a sentence fragment for subjects to complete (e.g. *Mary accidentally scratched John with a knife because_____*). Each fragment was completed by at least 32 (or as many as 45) subjects. The mean proportion of continuations that referred to the first character name (out of all continuations that referred to one or the other of the characters) was .46. The variability across items was high, but conditionalizing the response time data (given below) on the relative proportions of continuations did not yield any meaningful differences in the patterns of response times.

Procedure, Design, and Subjects. The procedure in Experiment 7 was the same as for Experiments 1 and 3. There were two variables in the design: The second clause of the final sentence used a pronoun intended to refer either to the first or second character mentioned in the first clause, and the test word was either the first character's name or the second character's name. These four conditions were combined in a Latin square with the 32 texts and 24 subjects (from the same population as the previous experiments).

Results

The data were analyzed in the same way as in the previous experiments, and the means are shown in Table 5. Response times longer than 2000 ms were eliminated (less than 1% of the data).

Insert Tables 5 and 6 about here

The main result is that there was no matching effect. Response time for a test word did not depend on whether the test word matched the intended referent of the pronoun that preceded it. Instead, response times were slower for the first character's name than the second character's name, whichever pronoun was used. This effect was significant, $F_1(1,23)=4.8$ and $F_2(1,31)=4.8$. Other F s for response times were less than 1.0. There were also more errors on the first character's name, $F_1(1,23)=8.1$ and $F_2(1,31)=4.3$. For errors, the interaction between the pronoun and test word variables approached significance in the subjects' analysis, $F_1(1,23)=3.6$ and was significant in the items' analysis, $F_2(1,31)=5.4$. The other F s for the errors analysis were less than 2.0.

It is worth repeating here that conditionalizing the response time data on the continuation data did not yield a meaningful pattern of results. Neither in this experiment nor in Experiments 5 and 6 could failures to find a matching effect be predicted from continuation probabilities. In Experiments 5 and 6, subjects were likely to continue a sentence containing a subject-initiating verb with a pronoun referring to the subject character, but there was no matching effect. They were not particularly likely to continue a sentence containing an object-initiating verb with a pronoun referring to the object, and there still was no matching effect. The implication of these results is that, while continuation data may sometimes be helpful in eliciting subjects' intuitions, it cannot take the place of other kinds of tests of comprehension.

General Discussion

We hypothesized that the lexical representation of verbs exhibiting implicit causality guides comprehension of sentences that use those verbs. The lexical representation of these verbs calls for arguments that satisfy the roles of initiator and reactor. The verbs express some action or change of state on the part of the reactor for which some action or property of the initiator is the cause. For some verbs, the initiator appears in the subject position in the surface

structure of a sentence, and the reactor in the object position; for others, the surface position of the roles is the reverse. In both cases, we hypothesized that the relative accessibility of the initiator in the discourse model constructed during reading is increased. Additionally, because the verbs express an action or state of affairs brought about by the initiator, it is natural that a *because* clause following the verb explain the initiator's behavior. The increased accessibility of the initiator, the natural fit of the explanation to the verb's lexical structure, and the use of the connective *because* together support pronoun resolution in sentences in which a verb exhibiting implicit causality is followed by an explanatory clause consistent with it. In the sentence *John blamed Mary because she forgot the wine*, the action of blaming is initiated by Mary (something she did), and the reason that she brought about blaming is that she forgot the wine. Mary is more accessible than John, it is natural to explain how she caused blaming, and *because* makes the causal relation explicit; these factors together support identification of Mary as the referent of the pronoun. In contrast, for the sentence *John blamed Mary because he was in such a bad mood*, the gender of the pronoun is not consistent with the more accessible of the two characters, and the explanation of the blaming action does not immediately fit with the implicit causal structure of the verb. These factors work against identification of Mary as the referent of the pronoun, and support the alternative referent, John.

Our data showed a matching effect both for pronouns consistent with a verb's causality and for pronouns inconsistent with it: Responses to a character's name were faster when the character was the referent of the pronoun than when it was not. There are at least two possible ways to describe the decision process that leads to this difference in response times. One possibility is that the test word is matched against the already existing representation of the sentence in memory, and response time and accuracy for the test word reflect its accessibility in that representation. In this case, the test word has no effect on the existing representation, and the information provided by the test word interacts with information in the text only in ways that produce no new information about the text. A second possibility is that the test word is used as

additional information in that it changes the text representation (Forster, 1981). In terms of our experiments, this could mean that the pronoun's referent had not yet been completely identified before the test word was presented, but that when the referent was presented as a test word, subjects at that point matched it against the pronoun and the discourse representation to identify it as the referent. (Note, though, that this identification process would have to act very quickly; there was only 200-300 ms between display of the last word of the sentence and presentation of the test word). Of course, presenting the referent as a test word does not add any really new information; the referent word is already in short-term memory because it was just mentioned in the preceding clause (Clark & Sengul, 1979). But presenting it as a test word could, for example, add to its accessibility sufficiently that pronoun resolution could succeed when it had not already. If correct, this second possibility would make the pronoun resolution that appears in our experiments critically dependent on the presence of the test word. In striking contrast, pronoun resolution in previous experiments (Greene et al., 1992, Experiments 1, 2, 3, 4, & 7) could not even have been dependent on the presence of a test word; in those experiments, there was no evidence that the referents of pronouns were identified at all.

In the experiments reported by Greene et al. (1992), we used sentences like *Mary accidentally scratched John with a knife and then she dropped it on the counter*. The main verbs in these sentences do not have implicit causality as a central part of their lexical representations (see Levin, 1992, for a discussion of *scratch*, for example). Therefore, we suggested, they do not privilege one of their arguments over the other. When discourse models are constructed during reading for sentences like these, the two arguments are not differentially accessible, and the second clause is not naturally attributed to one argument or the other by the structure of the verb. When a pronoun in the second clause is matched against the discourse model, the two arguments do not differ in accessibility, and the pronoun is not identified as referring to one or the other of them.

The results presented here suggest that one way a discourse can support pronoun resolution is by using a verb that increases the accessibility of one possible referent more than that of another and by attributing to the pronoun's referent information that fits naturally with the meaning of the verb. In these circumstances, and possibly in others, pronoun resolution may even be a mandatory component of comprehension (Gerrig, 1986). In contrast, as was the case with the materials used by Greene et al. (1992, Experiments 1-7), when a discourse does not support the identification of a unique referent for a pronoun, either because no referent is sufficiently accessible or because several possible referents are all equally accessible, then special goals or strategies may be required. In some of the experiments reported by Greene et al., the procedure was almost identical to that used in the experiments reported in this article: a reading speed normal for college undergraduates (Greene et al. used a constant 250 ms/word pace compared to the 170 ms/word plus 17 ms per letter we used), and subjects were given no specific task that required identification of pronominal referents. The data showed no evidence that unique referents for pronouns were identified. Evidence of pronoun resolution appeared only when test locations were made highly predictable by using just one-sentence texts, when subjects were motivated by a specific task that required pronoun resolution, and when they were given ample time to accomplish the resolution process during reading by presenting the words of the sentences at a rate of about 500 ms each.

As we and others have noted, in natural discourse, pronouns are typically used when only one entity is already highly salient in the comprehender's discourse model (Brennan, 1989; Chafe, 1974; Ehrlich, 1980; Fletcher, 1984; Greene et al., 1992). Use of verbs that exhibit implicit causality is only one of many ways in which natural discourse may make one entity more salient than others, and thereby support pronoun resolution. A variety of other devices may also be used to increase the accessibility of one entity: the cataphoric *this* ("This man walks into a bar...", Gernsbacher & Shroyer, 1989), cleft sentences ("It was Umberto who...", Sidner, 1983a), repetition of a full noun phrase ("Number thirty passes to forty-one. Forty-one shoots,

and he misses," Brennan, 1989), and spoken stress (Brennan, 1989). In short, natural discourse may be designed precisely so that pronoun resolution can be accomplished without requiring any specific strategy on the part of the comprehender. We discuss the process of pronoun resolution here, as in Greene et al., not in terms of what the pronoun does to trigger a search for its referent, but instead in terms of what the discourse does to make such a search unnecessary—how it introduces entities so as to make anaphoric reference felicitous.

More generally, these results and those of Greene et al. speak to the kinds of research needed in discourse comprehension. It has recently been proposed that the representation of discourse comprehenders construct without specific goals or strategies is "minimal" (McKoon & Ratcliff, 1992). A minimal representation does not include all the inferences necessary to construct a full, real-life like mental model of the situation described by a text. Instead, the only inferences constructed are those that are based on easily available knowledge or that are required to achieve coherence with information that is in the same local part of the text. For example, by this view, inferences about "what will happen next" in a story are inferred only if they can be based on well-known information. What will happen next to an actress who falls off a 14th story roof is not well known, and, data suggested, not explicitly inferred (McKoon & Ratcliff, 1986). But what will happen next when a seamstress gets out a needle and thread can be inferred on the basis of well-known associations (McKoon & Ratcliff, 1989a, 1989b). The minimalist approach runs counter to most current constructionist theories of discourse processing (van Dijk & Kintsch, 1983; Glenberg, Meyer, & Lindem, 1987; Morrow, Greenspan, & Bower, 1987; Trabasso & van den Broek, 1985) in obvious ways: Death would be an inference required for a lifelike mental model of the actress and her fate, yet the minimalist claim is that it is not explicitly inferred during reading. However, the finding that pronoun resolution processes may fail to identify a unique referent for a pronoun pushes the minimalist approach much further. After all, inferring that someone dies after falling from a 14th story roof might be viewed as quite a complicated inference—unlike a pronoun, which is often thought to be trivially understood

by a reader. Clearly, from the pattern of results shown in this article and by Greene et al., pronoun resolution is not a trivial matter. The unanticipated nature of this pattern of results reinforces the minimalist emphasis on the importance of examining the local representation of discourse during comprehension. This pattern of results also underscores the minimalist claim that readers do not necessarily comprehend a discourse in some full, completely correct, way; some sorts of "comprehension" may give only an incomplete representation of the meaning of a text.

Prior to this set of experiments, it would have been difficult to guess that stylistically appropriate pronouns were not always understood, that their comprehension depended on the verbs that preceded them in their discourse, and that their comprehension depended on the kind of clause in which they were placed. It would have seemed far fetched to claim that the lexical representations of a verb could determine whether or not a pronoun in a different clause was understood. Here, we have expressed only the first preliminary ideas about how local representations of discourse might be constructed and what kinds of information they might depend upon, and only the first preliminary data to address these problems. But these data should be sufficient to indicate how much we don't know about even the "smallest" parts of discourse comprehension.

References

- Au, T.K. (1986). A verb is worth a thousand words: The causes and consequences of interpersonal events implicit in language. *Journal of Memory and Language*, 25, 104-122.
- Brennan, S.E. (1989). *Centering attention in discourse*. Unpublished manuscript, Stanford University.
- Brown, R., & Fish, D. (1983). The psychological causality implicit in language. *Cognition*, 14, 237-273.
- Caramazza, A., Grober, E., Garvey, C., & Yates, J. (1977). Comprehension of anaphoric pronouns. *Journal of Verbal Learning and Verbal Behavior*, 16, 601-609.
- Chafe, W.L. (1974). Language and consciousness. *Language*, 50, 111-133.
- Chang, F.R. (1980). Active memory processes in visual sentence comprehension: Clause effects and pronominal reference. *Memory and Cognition*, 8, 58-64.
- Clark, H.H. (1977). Bridging. In P.N. Johnson-Laird & P.C. Wason (Eds.), *Thinking: Readings in cognitive science* (pp. 411-420). New York: Cambridge University Press.
- Clark, H.H., & Sengul, C.J. (1979). In search of referents for nouns and pronouns. *Memory and Cognition*, 7, 35-41.
- Corbett, A.T., & Chang, F.R. (1983). Pronoun disambiguation: Accessing potential antecedents. *Memory and Cognition*, 11, 283-294.
- Dell, G.S., McKoon, G., & Ratcliff, R. (1983). The activation of antecedent information during the processing of anaphoric reference in reading. *Journal of Verbal Learning and Verbal Behavior*, 22, 121-132.
- van Dijk, T.A., & Kintsch, W. (1983). *Strategies of discourse comprehension*. New York: Academic Press.
- Ehrlich, K. (1980). Comprehension of pronouns. *Quarterly Journal of Experimental*

Psychology, 32, 247-255.

- Fillmore, C.J. (1971). Verbs of judging: An exercise in semantic description. In C.J. Fillmore & D.T. Langendoen (Eds.), *Studies in linguistic semantics* (pp. 273-296). New York: Holt, Rinehart, & Winston.
- Fletcher, C. (1984). Markedness and topic continuity in discourse processing. *Journal of Verbal Learning and Verbal Behavior*, 23, 487-493.
- Fodor, J.A., & Bever, T.G. (1965). The psychological reality of linguistic segments. *Journal of Verbal Learning and Verbal Behavior*, 4, 414-420.
- Fodor, J.A., Garrett, M., & Bever, T.G. (1968). Some syntactic determinants of sentential complexity, II: Verb structure. *Perception and psychophysics*, 3, 453-461.
- Forster, K.I. (1981). Priming and the effects of sentence and lexical contexts on naming time: Evidence for autonomous lexical processing. *Quarterly Journal of Experimental Psychology*, 33, 465-495.
- Garvey, C., & Caramazza, A. (1974). Implicit causality in verbs. *Linguistic Inquiry*, 5, 459-464.
- Gernsbacher, M.A. (1989). Mechanisms that improve referential access. *Cognition*, 32, 99-156.
- Gernsbacher, M. A. (1990). *Language comprehension as structure building*. Hillsdale, NJ: Erlbaum.
- Gernsbacher, M.A., & Shroyer, S. (1989). The cataphoric use of the indefinite this in spoken narratives. *Memory & Cognition*, 17, 536-540.
- Gerrig, R.J. (1986). Process models and pragmatics. In N.E. Sharkey (Ed.), *Advances in Cognitive Science* (pp. 23-42). Chichester, England: Ellis Horwood.
- Givon, T. (1976). Topic, pronoun, and grammatical agreement. In C.N. Li (Ed.), *Subject and topic* (pp. 149-188). New York: Academic Press.
- Glenberg, A.M., Meyer, M., & Lindem, K. (1987). Mental models contribute to foregrounding during text comprehension. *Journal of Memory and Language*, 26, 69-83.
- Greene, S.B., McKoon, G., & Ratcliff, R. (1992). Pronoun resolution and discourse models.

Journal of Experimental Psychology: Learning, Memory, and Cognition, 18, 266-283.

- Grosz, B. (1981). Focusing and description in natural language dialogues. In A.K. Joshi, B.L. Webber, and I.A. Sag (Eds.), *Elements of discourse understanding* (84-105). Cambridge, MA: Cambridge University Press.
- Grosz, B.J., Joshi, A.K., & Weinstein, S. (1983). Providing a unified account of definite noun phrases in discourse. In *Proceedings of the 21st Annual Meeting of the Association of Computational Linguistics* (pp. 44-50).
- Grosz, B., & Sidner, C. (1986). Attention, intentions and the structure of discourse. *Computational Linguistics*, 12, 175-204.
- Healy, A.F., & Miller, G.A. (1971). The relative contribution of nouns and verbs to sentence acceptability and comprehensibility. *Psychonomic Science*, 21, 94-96.
- Hoffman, C., & Tchir, M.A. (1990). Interpersonal verbs and dispositional adjectives: The psychology of causality embodied in language. *Journal of Personality and Social Psychology*, 58, 765-778.
- Hudson, S.B., Tanenhaus, M.K., & Dell, G.S. (1986). The effect of the discourse center on the local coherence of a discourse. In *Proceedings of the Eighth Annual Conference of the Cognitive Science Society* (pp. 96-101).
- Just, M.A., & Carpenter, P.A. (1980). A theory of reading: From eye fixations to comprehension. *Psychological Review*, 87, 329-354.
- Levin, B. (1992). *English verb classes and alternations: A preliminary investigation*. Forthcoming, MIT Press.
- MacDonald, M.C., & MacWhinney, B. (1990). Measuring inhibition and facilitation from pronouns. *Journal of Memory and Language*, 29, 469-492.
- Matthews, A., & Chodorow, M. (1988). Pronoun resolution in two-clause sentences: Effects of ambiguity, antecedent location and depth of embedding. *Journal of Memory and Language*, 27, 245-260.

- McKoon, G., & Ratcliff, R. (1980). The comprehension processes and memory structures involved in anaphoric reference. *Journal of Verbal Learning and Verbal Behavior*, 19, 668-682.
- McKoon, G., & Ratcliff, R. (1984). Priming and on-line text comprehension. In D.E. Kieras & M.A. Just (Eds.), *New methods in reading comprehension research* (pp. 119-128). Hillsdale, NJ: Erlbaum.
- McKoon, G., & Ratcliff, R. (1986). Inferences about predictable events. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 12, 82-91.
- McKoon, G., & Ratcliff, R. (1989a). Inferences about contextually-defined categories. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 15, 1134-1146.
- McKoon, G., & Ratcliff, R. (1989b). Semantic association and elaborative inference. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 15, 326-338.
- McKoon, G., & Ratcliff, R. (1992). Inference during reading. *Psychological Review*, forthcoming.
- McKoon, G., Ratcliff, R., Ward, G., & Sproat, R. (1992). *Syntactic prominence effects on discourse processes*. Manuscript submitted for publication.
- McKoon, G., Ward, G., Ratcliff, R., & Sproat, R. (1992). Morphosyntactic and pragmatic factors affecting the accessibility of discourse entities. *Journal of Memory and Language*, forthcoming.
- Miller, G.A. (1962). Some psychological studies of grammar. *American Psychologist*, 17, 748-762.
- Morrow, D., Greenspan, S. & Bower, G. (1987). Accessibility and situation models in narrative comprehension. *Journal of Memory and Language*, 26, 165-187.
- Osgood, C.E. (1970). Interpersonal verbs and interpersonal behavior. In J.L. Cowan (Ed.), *Studies in thought and language* (pp. 133-228). Tucson, AZ: University of Arizona Press.
- Rayner, K. (1978). Eye movements in reading and information processing. *Psychological*

Bulletin, 85, 618-660.

Sidner, C.L. (1983a). Focusing in the comprehension of definite anaphora. In M. Brady & R. Berwick (Eds.), *Computational models of discourse* (pp. 267-330). Cambridge, MA: MIT Press.

Sidner, C.L. (1983b). Focusing and discourse. *Discourse Processes*, 6, 107-130.

Trabasso, T., & van den Broek, P. (1985). Causal thinking and the representation of narrative events. *Journal of Memory and Language*, 24, 612-630.

Ward, G., Sproat, R., & McKoon, G. (1991). A pragmatic analysis of so-called anaphoric islands. *Language*, 67, 439-474.

Webber, B. (1983). So what can we talk about now? In M. Brady and R. Berwick, (Eds.), *Computational models of discourse* (pp. 331-371). Cambridge, MA: MIT Press.

Author Note

This research was supported by NIDCD grant R01-DC01240 and AFOSR grant 90-0246 (jointly funded by NSF) to Gail McKoon and by NIMH grants HD MH44640 and MH00871 to Roger Ratcliff. We thank Beth Levin for discussions of this work. Correspondence concerning the article should be addressed to Gail McKoon, Psychology Department, Northwestern University, Evanston, IL, 60208.

Table 1
Examples of Experimental Texts

Subject-initiating verbs

James and Debbie were working on a political campaign together.

They were both planning on pursuing careers in politics.

James infuriated Debbie because

he leaked important information to the press.

she had to write all the speeches.

Object-initiating verbs

The boss had been giving Diane and Sam a hard time lately.

Finally the two of them decided to do something about it.

Diane valued Sam because

he always knew how to negotiate.

she never knew how to negotiate.

Table 2

Results of Experiments 1 through 4: Response Times and Error Rates

	RT	% Errors	RT	% Errors
Subject-initiating verbs	Experiment 1		Experiment 2	
<u>Test first character</u>				
Consistent continuation				
(Referent matches test)	1005	5	776	7
Inconsistent continuation				
(Referent doesn't match test)	1083	0	780	5
<u>Test second character</u>				
Consistent continuation				
(Referent doesn't match test)	1130	2	835	6
Inconsistent continuation				
(Referent matches test)	1060	2	795	4
Object-initiating verbs	Experiment 3		Experiment 4	
<u>Test second character</u>				
Consistent continuation				
(Referent matches test)	933	2	733	5
Inconsistent continuation				
(Referent doesn't match test)	974	1	764	4
<u>Test first character</u>				
Consistent continuation				
(Referent doesn't match test)	1008	9	784	12
Inconsistent continuation				
(Referent matches test)	957	3	735	5

Table 3

Results of Experiments 5 and 6: Response Times and Error Rates

	RT	% Errors
Subject-initiating verbs	Experiment 5	
<u>Test first character</u>		
Consistent continuation		
(Referent matches test)	934	2
Inconsistent continuation		
(Referent doesn't match test)	918	1
<u>Test second character</u>		
Consistent continuation		
(Referent doesn't match test)	921	2
Inconsistent continuation		
(Referent matches test)	917	1
Object-initiating verbs	Experiment 6	
<u>Test second character</u>		
Consistent continuation		
(Referent matches test)	880	3
Inconsistent continuation		
(Referent doesn't match test)	887	3
<u>Test first character</u>		
Consistent continuation		
(Referent doesn't match test)	938	5
Inconsistent continuation		
(Referent matches test)	951	5

Table 4

Example of Paragraphs from Experiment 7

Mary and John were doing the dishes after dinner.

One of them was washing while the other dried.

Mary accidentally scratched John with a knife because

she was so tired and clumsy.

he suddenly grabbed for a glass.

Table 5

Results of Experiment 7: Response Times and Error Rates

	RT	% Errors
<u>Test first character</u>		
Referent matches test	975	5
Referent doesn't match test	978	2
<u>Test second character</u>		
Referent doesn't match test	947	1
Referent matches test	930	2

Table 6
Response Times and Error Rates for Filler Test Words and True/False Test Sentences,
and Standard Errors of the Means

Exp	Pos. Test Wds.		Neg. Test Wds.		True Test Sts.		False Test Sts.		S.E.
	Err.		Err.		Err.		Err.		
	RT	%	RT	%	RT	%	RT	%	
1	1253	12	1237	4	2437	7	2259	12	22
2	932	24	890	9	2086	8	2060	12	10
3	1141	16	1094	3	2240	9	2181	15	19
4	888	22	857	9	1982	8	1987	17	17
5	1071	13	1028	8	2162	7	2076	13	18
6	1083	14	1030	5	1999	8	1987	15	16
7	1128	16	1074	4	2050	8	1962	12	14

Note: S.E. refers to the error in the means of the experimental conditions tested by ANOVA.