

Listeners as Co-Narrators

Janet B. Bavelas, Linda Coates, and Trudy Johnson
University of Victoria

A collaborative theory of narrative story-telling was tested in two experiments that examined what listeners do and their effect on the narrator. In 63 unacquainted dyads (81 women and 45 men), a narrator told his or her own close-call story. The listeners made 2 different kinds of listener responses: *Generic* responses included nodding and vocalizations such as "mhm." *Specific* responses, such as wincing or exclaiming, were tightly connected to (and served to illustrate) what the narrator was saying at the moment. In experimental conditions that distracted listeners from the narrative content, listeners made fewer responses, especially specific ones, and the narrators also told their stories significantly less well, particularly at what should have been the dramatic ending. Thus, listeners were co-narrators both through their own specific responses, which helped illustrate the story, and in their apparent effect on the narrator's performance. The results demonstrate the importance of moment-by-moment collaboration in face-to-face dialogue.

These experiments are part of a larger program of research on the social nature of language use in face-to-face dialogue (e.g., Bavelas, 1990; Bavelas, Hutchinson, Kenwood, & Matheson, 1997; Watzlawick, Beavin Bavelas, & Jackson, 1967). We propose that face-to-face dialogue is shaped by social as well as syntactic and semantic processes. That is, dialogue is more than the individuals' production and comprehension of language; there are essential on-line collaborative processes as well. The nature and importance of these collaborative processes has been demonstrated in experiments in which both participants had a speaking role (e.g., Bavelas, Chovil, Coates, & Roe, 1995; Clark & Wilkes-Gibbs, 1986; Schober & Clark, 1989). In the present research, we studied asymmetrical dialogues, in which one person was telling a story and the other person was "merely listening." That is, we asked what an ostensibly passive listener does and what effect he or she can have on the narrator.

The Listener in Language and Communication Theories

Listeners have at best a tenuous foothold in most theories. At the extreme, listeners are considered nonexistent or irrelevant because

the theory either does not mention them or treats them as peripheral. This omission may be attributed, in part, to the implicit use of written text as the prototype for all language use (Linell, 1982), so that listeners are functionally equated with readers (e.g., Kintsch & van Dijk, 1978, p. 364). If the writer has an imagined reader in mind, he or she may shape the writing accordingly. When this model is extended to conversation, the imagined listener has a similarly abstract role. The speaker's perception of the listener becomes the hypothetical target: "In essence, the hearer is dealt with only as a figment of the speaker's imagination and not as an active coparticipant in [his or her] own right." (Goodwin, 1986, p. 205). Because there is no mention of any moment-by-moment influence in the course of the actual conversation, the listener remains "mute or invisible" (Clark & Wilkes-Gibbs, 1986, p. 3).

A somewhat less monologic view treats the listener as a "speaker-in-waiting." That is, the listener is present but not active during the other's speech; he or she is simply awaiting the speaking turn. This approach handles the differences between monologue and dialogue by casting conversation as alternating monologues. Once this structure is imposed, the focus remains on the soliloquy of the person who has the floor, with the listener as a passive audience. (Indeed, the term *floor* derives from formal legislative debate and refers to "the right of one member to speak . . . in preference to other members;" *Random House Unabridged Dictionary*, 2nd ed.) Thus, the considerable research interest in conversational turn-taking, especially smooth or successful turn-taking, can be traced to the premise that there exists a conversational floor that can only properly be held by one person at a time. Listener responses that do not claim the floor are treated as problematic or are ignored.

In addition to models based on written language and formal monologues, another strong theoretical influence that excludes listener activities is the classic Shannon and Weaver (1949) model of information transmission, introduced to psychology by Miller (1951). In this model, there is only a one-way channel from sender to receiver at any given moment; the receiver has no way to respond until he or she becomes the sender and takes over the channel. The classic model is deeply embedded in the terms we

Janet B. Bavelas, Linda Coates, and Trudy Johnson, Department of Psychology, University of Victoria, British Columbia, Canada.

Linda Coates is now at Okanagan University College, Kelowna, British Columbia, Canada; and Trudy Johnson is now with R. A. Malatest & Associates, Ltd., Victoria, British Columbia, Canada.

These studies were supported by grants from the Social Sciences and Humanities Research Council of Canada. Earlier versions of these data were presented to the Stanford Language Use Group, May 1994; at the Canadian Psychological Association Annual Meeting, June 1994; and at the International Communication Association Annual Meeting, May 1995. We would like to thank our analysts and raters, especially Jeffrey Hancock, Patricia Hart, Catherine Noeth, and Gillian Roberts.

Correspondence concerning this article should be addressed to Janet B. Bavelas at the Department of Psychology, University of Victoria, P.O. Box 3050, Victoria, British Columbia V8W 3P5, Canada. Electronic mail may be sent to jbb@uvic.ca.

still use to describe conversational participants and processes, such as *sender*, *receiver*, and *channel*.

Yngve (1970) pointed out that these simple distinctions do not apply to actual conversation:

[T]he distinction between having the turn or not is not the same as the traditional distinction between speaker and listener, for it is possible to speak out of turn, and it is even reasonably frequent that a conversationalist speaks out of turn. In fact, both the person who has the turn and [his or her] partner are simultaneously engaged in both speaking and listening. This is because of the existence of what I call the back channel, over which the person who has the turn receives short messages such as "yes" and "uh-huh" without relinquishing the turn. (p. 568)

On the basis of his observations, Yngve modified the classic model by creating a parallel, although definitely subordinate, *back channel*: The speaker owns the main channel, and the listener makes minimal, noninterruptive responses on the back channel. In one interpretation of this model, the listener's responses are simply shunted off to the back channel, where they would have no effect on the speaker. If no connection is made between back channels and the narrator's behavior, the implication is that back channels are merely appropriate behavior in the listener role and are (effectively) emitted randomly.

In brief, all of the above approaches fit what Schober and Clark (1989) called the *autonomous* view of conversation, in which the speaker delivers information in polished monologues for a passive listener to comprehend. There is no specification of an active role for the listener in the conversation. However, Yngve proposed some reciprocal effect, namely, that "the back channel appears to be very important in providing for monitoring the quality of communication" (p. 568). Indeed, experimental research into the effect of feedback on communication has a longer history, beginning with Leavitt and Mueller (1951). Subsequently, a few social psychologists have demonstrated listener effects experimentally, showing that limiting or removing listener responses affected the efficiency or effectiveness of the speaker's encoding. Krauss and Weinheimer (1966) found that speakers giving instructions over an intercom used more words when listener feedback was reduced or eliminated. Using a similar task, Krauss, Garlock, Bricker, and McMahon (1977) found that delayed feedback from the listener also increased the number of words the speaker used to send the information. Kraut, Lewis, and Swezey (1982) asked speakers to summarize a movie with varying amounts of listener feedback; the listeners understood better when they could provide feedback.

A model that fits these data is Clark's (1996) *collaborative* model, in which "speakers and their addressees go beyond . . . autonomous actions and collaborate with each other moment by moment to try to ensure that what is said is also understood" (Schober & Clark, 1989, p. 211). In other words, language in dialogue is a joint activity. The actions that make up a dialogue are not engaged in independently but rather require constant coordination; dialogue is a duet, not two solos (Clark, 1996, especially chapters 2 & 3). Other scholars, most notably the conversation analysts, have also proposed that discourse is a joint activity (e.g., Atkinson & Heritage, 1984; Goodwin, 1979, 1981; Goodwin & Goodwin, 1987; Sacks, 1974; Sacks, Schegloff, & Jefferson, 1974; Streeck, 1994).

Clark and Wilkes-Gibbs (1986) and Schober and Clark (1989) have experimentally demonstrated the moment-by-moment collaboration and coordination of dialogue. They used a task in which the speaker had to describe which figure in a group of unusual figures the addressee should choose next. Even though it was the speaker who had the answer and the task of finding a suitable reference for each figure, the addressee in fact participated actively, not only by conveying understanding (or lack of understanding), but also by offering candidate descriptions that were often adopted by the speaker. In the following example, S is the speaker and A the addressee:

- S: Then [figure] number 12 is (laughs) looks like a, a dancer or something really weird. Um, and, has a square head, and um, there's like there's uh, this kinda this um,
 A: Which way is the head tilted?
 S: The head is, eh, towards the left, and then the, an arm could be like up towards the right?
 A: Mm-hm.
 S: And, it's-
 A: [overlapping] an, a *big fat leg*? You know that one?
 S: [overlapping] *Yeah, a big fat leg.*
 A: and a little leg.
 S: Right.
 A: Okay.
 S: Okay?
 A: Right. (Adapted from Schober & Clark, 1989, pp. 216–217; italics added)

In later trials, the speaker incorporated the addressee's description, referring to this figure as "the dancer with the big fat leg." Thus, although in the autonomous view the linguistic activity of reference or referring (i.e., finding the right word) is traditionally the speaker's prerogative and responsibility, the final description was a joint product, tailored uniquely by speaker and addressee working together. As Clark and Wilkes-Gibbs (1986) pointed out, referring was a collaborative process.

We are interested in extending the collaborative model and previous research on listener effects in two ways. First, in the studies just described (Clark & Wilkes-Gibbs, 1986; Schober & Clark, 1989), the addressee was arguably an active participant because he or she did have some of the information needed (i.e., the array of possible figures) and could therefore contribute by making suggestions. Equally important, the addressee explicitly needed to get the information right for both of them to succeed at the task; these were inherently collaborative tasks. Examining the possibility of a truly passive listener requires that the addressee have none of the speaker's information and no formal collaborative role, as in the task we used here: The speaker told a stranger about a close call he or she had had in the past. These addressees should more closely approximate listeners as portrayed in traditional, autonomous theories, because they have no information to contribute and no formal role to play in the story-telling. However, if referring is a collaborative process, perhaps narrating is as well:

Narratives seem different from conversations, because they seem to be produced by individuals speaking on their own. . . . but appearances belie reality. Narratives rely just as heavily on coordination among the participants as conversations do. It is simply that the coordination is hidden from view. (Clark, 1994, pp. 1006–1007)

We are interested in bringing this coordination out of hiding by examining closely what listeners do during narration.

A second change we have introduced is to study listeners in face-to-face dialogue, rather than, for example, on intercom links or with a partition between participants, as in previous research. To understand why this difference might be relevant, it is necessary to ask what is unique about face-to-face dialogue. Fillmore (1981), Linell (1982), and Clark (1996) have proposed that face-to-face dialogue is a "basic" or "primary" setting for language use, both because of its ubiquity and because it has a number of unique features not found in other settings. Two of the features explicated by Clark (1996, pp. 8–10) are *visibility* and *simultaneity*. In face-to-face dialogue, the participants can see each other and they "can produce and receive at once and simultaneously" (p. 9). These two features, combined, have particular significance for listeners. Because the speaker can see as well as hear their responses, listeners do not need to interrupt vocally; they have available a much wider repertoire of simultaneous but noninterruptive responses, especially facial displays such as nodding, smiling, looking confused, or wincing (Bavelas & Chovil, 1997; Chovil, 1991/1992), than they would in any other setting. Goodwin (1981) showed that speakers are highly sensitive to listener gaze; if they start a sentence and discover the listener is not looking at them, they restart (and often rephrase) when the listener looks back. Therefore, listeners can arguably contribute more fully in face-to-face dialogue than when their visible reactions are not seen.¹ (We do not mean to imply that there will always be differences between face-to-face and mediated settings, because of the particular demands of different tasks and the likelihood of adaptation. That is, participants may not always need visibility or simultaneity for their task, or they may find other ways of accomplishing these functions; see Phillips & Bavelas, 2000; Williams, 1977.)

Two Kinds Of Listener Responses

If, as stated above, our goal is to uncover the "hidden" ways in which listeners might be actively involved in the narrative process, we must begin with a closer examination of what listeners do. For the most part, listener responses (also called back channels) have been treated as a uniform class, the prototype of which are acts such as nodding and generic vocalizations (e.g., "mhm," "uh-huh," or "yeah"), which do not convey any narrative content. Lacking such content, they seem to function solely as an indicator of the listener's cognitive processes, which the narrator can use to track comprehension and make corrections if necessary. We propose the term *generic* listener responses to describe these standard examples of back channels. Generic listener responses are not specifically connected to what the narrator is saying, in the sense that the same generic response would be appropriate to a wide variety of narratives. For example, one might appropriately nod while listening to a lecture, a sad story, or an exciting story. Here are two typical examples of generic listener responses from our close-call stories; the listener's response is placed exactly below the point it occurred in relation to the narrator's utterance:

- Example 1
Narrator: "We stayed in an RV park."
Listener: ["Mhm" with nod]
- Example 2
Narrator: "I have a single bed . . . with a headboard?"
Listener: [nod] [nod with "mm"]

Notice that these responses are timed precisely and appropriately, but they are not specific to the narrative content of the moment.

Most close observers (e.g., Goodwin, 1986; Krauss et al., 1977; Kraut et al., 1982; Yngve, 1970) have also described what seems to us a different kind of response. Yngve gave the following excerpt from his data:

He says, "When you've accumulated possessions . . .," and she says, "*piece-by-piece*," simultaneously with his uttering the word "possessions." It is not a case of her attempting to supply a word he can't think of, for there is no hesitation on his part. Her simultaneous activity can be analyzed as an example of her agreeing with what he is saying by volunteering appropriate words instead of mere indications of assent, such as "yes." This analysis is supported by the considerable extent of [their] previous discussion on the accumulation of possessions, allowing her to predict easily what he might say. Her utterance has the intonational and gestural characteristics of enthusiastic and animated agreement. (p. 574; italics added)

We would also point out that, in addition to indicating agreement, the listener contributed an appropriate and specific description of the particular manner of accumulating possessions (i.e., "*piece-by-piece*" rather than all at once); that is, she helped him tell this part of the story.

We call these contributions *specific* listener responses; examples include looking sad, gasping in horror, mirroring the speaker's gesture, or supplying an appropriate phrase (as above). They include what has traditionally been called *motor mimicry* (such as wincing at another person's injury; Bavelas, Black, Lemery, & Mullett, 1986; Bavelas & Chovil, 1997), some hand gestures, as well as brief verbal interjections (see the *Analysis* section of Experiment 1, below, for operational definitions.) Specific listener responses are tightly connected to what the narrator is saying at the moment. They take on a form specific to the narrative content of the moment and are not generically appropriate to all narratives. Our data included these examples:

- Example 3
Narrator: He flipped his truck over, over an embankment.
Listener: [facial display of concern]

The listener's facial display is specific to the incident being described. Once the truck had flipped over (and not before), the narrator would have been concerned.

- Example 4
Narrator: No one was around, and he said, 'Get in the car.'
Listener: [facial display of fear]

Similarly, in this example, when the stranger suddenly ordered the narrator into his car, she would have been afraid. Notice that, in both examples, the speaker could easily have made these facial displays to illustrate his or her own narrative. However, in both cases, the speaker only described the facts; it was the listener who illustrated the dramatic or emotional import of the facts (concern or fear) and thereby enriched the story.

¹ Listeners also recall less when there is not a "live" speaker, for example, hearing the same sentence on a tape recording versus in person (Feldman, 1971), which may be due to simultaneity or visibility or both.

Example 5

Narrator: I, like an idiot, decide to climb up the cliff instead of . . .

Listener: . . . going up the road

Narrator: . . . taking the easy way out and going up the road.

Here, the listener actually provided a phrase that fit the narrator's main point, which was that it was foolish to try to escape rising water by going up the nearby cliff when there had been a better option. Moreover, he did so in a way that exactly fit the narrator's syntax, and the narrator immediately incorporated the listener's interjection into the narrative.

Thus, specific responses permit listeners to become, for the moment, co-narrators who illustrate or add to the story. To do so, they must track the narrative very closely. In the close-call stories told in our experiments, narrators often moved quickly between horror and humor, both dramatizing the potential danger and making fun of it (because in the end there was no harm). The listeners' specific responses had to—and did—follow these rapid shifts. What is remarkable about the listeners' ability to coordinate so well is the fact that the narrators and listeners were strangers to each other, so the listeners had no previous knowledge of the story. Yet they were able to contribute specific and appropriate details, moment-by-moment. In summary, the distinction between generic and specific responses is a functional one, based on the relationship of the response to the narrative (and not on the listener's response in isolation). Table 1 summarizes several contrasts between generic and specific responses, along with important similarities.

It is important to point out that our generic-specific distinction does not map onto verbal-nonverbal. Nods are usually generic, whereas winces are specific, although both are nonverbal. The verbal *yeah* is usually a generic response but *going up the road* (in Example 5) is a specific response. More important, many listener

responses combine both audible and visible elements; for example, a nod plus *yeah* (generic) or a wince with *oh no!* (specific). Our *integrated message model* (Bavelas, 1994; Bavelas, Black, Chovil, & Mullett, 1990, chapter 6; Bavelas & Chovil, 1997, 2000) proposes that these audible and visible acts combine with each other moment-by-moment in the conversation to produce an integrated message (what Clark [1996, chapter 6, especially pp. 185–187] and Engle & Clark [1995] have called a *composite*). In these integrated messages, both audible and visible acts can be used to communicate the meaning. That is, we treat certain visible acts such as facial and hand gestures as part of the speaker's or listener's message, extending what Bruner (1990) called "acts of meaning" to include both audible and visible acts in face-to-face dialogue. Accordingly, our distinction between generic and specific listener responses depends on the meaning of the response in its narrative context and not on the physical channel in which it happens to be encoded. (Because our meaning-based approach depends on interpretation, interanalyst reliability will be crucial.)

As noted above, earlier researchers have described what we are calling specific listener responses but have not treated them as significantly different from generic listener responses. Goodwin (1986) is the notable exception. In his field studies, he made a functional distinction between two kinds of listener responses, which he called *continuers* and *assessments*:

Many of these vocalizations—the prototypic example being "uh huh"—function as "continuers," actions displaying [the] recipient's understanding that an extended turn at talk is in progress but not yet complete. . . . Some of the brief responses produced by recipients seem to go beyond this. . . . [Such a response], rather than simply acknowledging receipt of the talk just heard, assesses what was said by treating it as something remarkable. (p. 207)

Goodwin (1986) went on to posit that,

assessments display an analysis of the *particulars* of what is being talked about [and permit the] recipient to react to the talk in progress by showing enthusiasm, appreciation, outrage, et cetera. (p. 210; italics added)

Goodwin was interested in where these two different kinds of responses occurred in relation to the narrator's speech (between or within clauses). Our focus is on how these responses contribute to the narrative.

Summary of Predictions

The goal of the present research was to demonstrate experimentally the validity of the distinction between generic and specific listener responses, that is, to show that it is possible, necessary, and fruitful to distinguish between them. First, as noted above, the difference depends on an interpretation of the meaning of the response in relation to the narrative, rather than on any physical parameter. Therefore, it is essential to demonstrate that this interpretation can be rendered objective by achieving good agreement between independent analysts who are naive to our hypotheses.

Second, we agree with Goodwin (1986, p. 215) that specific responses should occur later in the narrative than do generic responses—at least for these kinds of stories. Many scholars (Fivush, 1991; Labov, 1972; Mandler, 1987; Neisser, 1982; Polanyi, 1989; Rumelhart, 1975) have pointed out that personal narratives

Table 1
Similarities and Differences Between Generic and Specific Listener Responses

Similarities	
Both are appropriate responses, related to the narrative.	
Both show that the listener is understanding, attending, following.	
Both occur at appropriate places within the narrative.	
Both may be responses to main points of the narrative or to digressions.	
Differences	
Generic responses	Specific responses
Are <i>listening</i> .	Are <i>co-telling</i> .
Keep the listener as <i>audience</i> or <i>observer</i> .	Make the listener an <i>actor</i> in the story.
Are made <i>to</i> or <i>at</i> the story or narrator.	Are made <i>with</i> the story or narrator.
Are <i>generally</i> related to the narrative.	Are <i>specific</i> to this point in the narrative.
Are <i>external</i> to the narrative plot.	Are <i>internal</i> to the narrative plot.
<i>Respond</i> to the narrative plot.	<i>Act upon</i> (add to) the narrative plot.
Communicate <i>general</i> understanding.	Communicate <i>specific</i> understanding.
Indicate understanding of the <i>words</i> .	Indicate understanding of the <i>implications</i> of the words.

have a conventional form. The typical close-call story told to a stranger requires a set-up phase to give the background; the story then proceeds to its dramatic conclusion. Generic responses are possible and appropriate as soon as the narrator begins, and they remain appropriate throughout the story. In contrast, specific responses require more information about the story and about the narrator's perspective (e.g., humor or horror). They would not be possible or credible until the listener has this information. Also, if we are correct that specific responses function to illustrate and co-narrate, they should be more appropriate later, during more dramatic parts of the narrative, where they can illustrate the narrator's main theme (danger, fear, etc.). We therefore predicted that specific responses would have a later serial position, on average, than generic responses.

Third, we created experimental conditions to examine the effects of involvement versus distraction on the listener's ability to respond. We predicted that, when the listener is distracted from the dialogue, he or she should be less able to make appropriate listening responses. Because specific responses have a closer relation to the narrative, they should be more affected by distraction.

Finally, we predicted that if dialogue (including storytelling) is always collaborative, then distracting the listener from the narration should affect the quality of the storytelling. That is, the narrator needs a listener to tell a good story; a good listener is a collaborator, a partner in storytelling. We tested these predictions in the following two experiments.

Experiment 1

Method

Participants

Eighty-six students recruited from the University of Victoria Psychology Department's volunteer subject pool were combined to form 43 dyads. All were strangers to each other. Three dyads were replaced and one was dropped for technical or procedural reasons, leaving 39 dyads for analysis. In these dyads, there were 57 women and 21 men; all dyads except one were the same gender. All of the participants consented to being videotaped in the Psychology Department's Human Interaction Laboratory and, after viewing their tape, gave permission for its subsequent analysis.

Equipment

The Human Interaction Lab has four remotely controlled Panasonic WD-D5000 color cameras and two special effects generators (a Panasonic WJ-5500B overlaid on a customized Panasonic four-camera system). We used two cameras to videotape both the narrator and the listener in a split-screen layout that recorded a face-on view of each of the participants and a time signal on the tape in minutes, seconds, and hundredths of seconds. For analysis, we used either a Sharp 2500S or JVC BR-S605-UB VHS VCR and a 19-inch (48 cm) Sony or Electrohome color monitor.

Procedure

The dyads were randomly assigned to one of four conditions (listening, summarizing, retelling, and counting), which will be explained in more detail below. In all four conditions, the experimenter began as follows:

What I'd like both of you to do is tell something about a close call or near-miss incident. A close call is something that happened where someone was almost hurt, or something bad almost happened, but in

the end everything turned out okay. Make sure that you tell something you're comfortable telling. And if you can't think of something that happened to you, then you can tell about a close call that happened to a friend. Just to give you some ideas, other people have told stories about skiing accidents, horseback-riding accidents, and nearly losing a paper on the computer. I would like you to tell your story in as much detail as you can. So don't just describe it in a couple of sentences.

The flip of a coin determined who told his or her story first, and the rest of the instructions were addressed to the listener.

The first three conditions varied in the degree of cognitive demand that the listener's assigned task imposed on him or her. In the listening condition ($n = 10$), the listener was to "just listen while [the narrator] tells [his or her] close-call story." In the summarizing condition ($n = 10$), the listener was to be prepared to, "in a sentence or less, very briefly summarize the gist of [the narrator's] story." In the retelling condition ($n = 9$), the listener had to be prepared to "retell [the narrator's] close-call story with as much detail as possible . . ."

The purpose of the fourth condition (counting, $n = 10$) was to prevent the listeners from listening to the narrator's story, that is, to distract them from the social interaction so that they could not pay attention to the story content. They were told,

[t]his may sound kind of bizarre, but what I'd like you to do is to count the number of days it is from now [in February] until Christmas. And if you get done with that, I would like you to count the number of statutory holidays from now until Christmas. And I'd like you to do this while [the narrator] is telling [his or her] close-call story.

In all conditions, both participants heard the full instructions together, before they took turns telling their personal close-call stories or being the listener. We analyzed the story told by the second narrator, when both participants were more relaxed.

Analyses

Listener responses. We analyzed listener responses made during the first minute of the narratives. First, we identified all of the listener responses in this segment, and transcribed them in relation to the words spoken by the narrator and the time of onset. Listener responses were defined for the analysts as,

Actions [both verbal and nonverbal] that indicate the person is attending, following, appreciating, or reacting to the story. They include (but are not limited to) nodding, "mhm," "yeah," smiling, laughing, motor mimicry, gesturing the content of the story, supplying words or phrases, dramatic intakes of breath, and displays of excitement, fear, or alarm.

At this point, we excluded noncommunicative behaviors made by the listener, such as adaptors, aborted communicative acts (e.g., the beginning of a smile that became a lip-bite), and behaviors that were too ambiguous to analyze.

Often, the listener's meaning was conveyed by several behaviors in concert. In these cases, we treated all of the behaviors that acted together to convey the same meaning as a single listener response. However, if the adjacent behaviors conveyed a different meaning, they were treated as separate responses. For example, a confused look that became a horrified expression would be two different listener responses. The reliability of this "packaging" procedure for two independent analysts was above 90%.

Using the system of analysis we had developed, two new, naive analysts then distinguished between what we called only "A" (generic) and "B" (specific) listener responses, as described in the following instructions:

"A" responses keep the listener clearly in the role of listener/audience. That is, they are made by the listener as *audience* or *observer* and are

external to the narrative plot. Like most listener responses, "A" responses are usually related to story content, but they are *not* part of the story plot. The listener selects out information in the speaker's narrative and responds to it. The listener is "removed" from the story plot (i.e., he or she is not critically involved in its creation). "A" responses are merely responses to the story that the listener makes in his or her role as audience. They include mostly what have been called back channels, and they tend to take standard forms, for example, nodding or saying "mhm" or "yeah." (Occasionally, a listener may make a response that is not related to the story content but is a comment on the communicative situation itself. These are also considered "A" responses.)

The specific listener responses were defined as follows:

"B" responses are *co-telling* acts in which the listener becomes (however briefly) involved *with* the narrator in the telling of the story. The listener is more of an actor than an observer of the story. "B" responses are *internal* to the narrative plot. The listener selects out information from the narrative plot and *acts upon* it. That is, the listener acts like someone in the story (or like the narrator while telling the story). In this way, the listener may be more involved with the creation of the story. "B" responses include such behaviors as motor mimicry, gesturing the content of the story, and supplying words or phrases that advance the content of the story.

We did not analyze smiling or laughing when they occurred alone, without any other communicative acts. We found no reliable way to distinguish smiles or laughs according to meaning although they did function as listener responses. Smiling and laughing could be polite or appreciative generic responses; they could also be specific to the narrator's own amusement, or they could be maintaining the dialogue (metacommunicative).

After learning to identify each listener response as an "A" or a "B", the two naive analysts applied the system of analysis, which included a formal decision tree,² to the first minute of each of the 39 dyads. They worked independently and were unaware of the listeners' condition, of the existence of different experimental conditions, and of our hypotheses. Reliability checks at the beginning, middle, and end of the analysis process revealed an overall agreement of 95%.

Story ratings. Other raters who were also unaware of condition rated the 39 stories on two separate parameters: how good the story plot was and how well the story was told. First, we summarized all of the plots briefly in point form. For example,

Narrator was a child, traveling by train with his family in London. Their first train was late, and then they got lost in the Underground. They were late for their connection and had to take another train. The train they missed was in an accident; 200 people died.

Seven undergraduates rated each synopsis on a scale ranging from 1 (*very little story potential*) to 5 (*very good story potential*). Their reliability was .69 (intraclass correlation); we used their averaged ratings.

Three new raters then viewed the tapes (with the listener's half of the screen covered) and rated each on a 5-point scale for how well it was told. The scale ranged from 1 (*very poor, even for an ordinary conversation*) to 5 (*excellent, for a nonprofessional*). Intraclass correlation was .83; we used their averaged ratings.

Results

Serial Position of Specific and Generic Responses

Because of the experimental effect (see below), specific responses disappeared entirely in some dyads, so we analyzed only the 25 dyads in which both specific and generic responses oc-

curred. We calculated the average serial position of specific and generic responses made within the first minute of the narrative for each dyad. As predicted, we found that the mean serial position of specific listener responses was significantly later in the narrative ($M = 11.42$, $SD = 4.94$) than was the position of generic responses ($M = 7.92$, $SD = 3.72$); $t(24) = 5.22$, $p < .0001$.

Effect of Condition on Rate of Generic and Specific Responses

There were no significant differences between the first three (attending) conditions, so we grouped them together for our analyses (see Figure 1). The attending listeners responded at high rates, making one kind of listener response or the other about every 3.5 s (generic: $M = 13.31$, $SD = 5.48$; specific: $M = 3.76$, $SD = 3.47$). In contrast, the counting condition affected the rate of both generic and specific responses. The listeners in the counting condition made significantly fewer generic responses than did listeners in the three attending conditions; $M = 6.10$, $SD = 4.25$; $t(37) = 3.78$, $p < .001$. The listeners who were counting also made significantly fewer specific responses than did listeners who were attending to the story; $M = .30$, $SD = .48$; $t(37) = 3.11$, $p < .004$.

Notice that the rate of generic responses in the counting condition was slightly less than half that of the attending conditions (6.1 vs. 13.3), whereas the rate of specific responses in the counting condition was less than one-tenth the rate in the attending conditions (0.30 vs. 3.73). We compared the differential effect of condition on response type in two ways, but neither supported this apparent difference: The effect sizes were very similar ($\eta^2 = .29$ and $.21$, respectively), and there was no significant interaction between experimental condition and generic versus specific response (treated as a second factor).

Effect of Experimental Condition on Quality of Narration

As we had hoped, there was no condition difference in the plot potential of the stories individuals decided to tell; $t(37) = .64$, $p = .53$. However, as predicted, narrators in the counting condition told their stories significantly less well than those in the attending condition, $M = 2.15$ ($SD = .86$) versus $M = 3.10$ ($SD = .98$); $t(37) = 2.72$, $p = .01$.

Other Variables: Gender and Story Length

We examined gender differences in the rate of generic, specific, and total listener responses. As we expected (Coates & Johnson, in press; Marche & Peterson, 1993), there were no significant differences; $t_s(37) = 1.31$, 1.58 , and 1.70 ; all p s $> .09$. Leaving aside one story that was almost 15 min long, the mean story length was 1 min 44 s ($SD = 48$ s). There was no effect of condition on story length, $t(36) = 0.27$, $p = .79$, and story length was not correlated with the storytelling ratings, $r = .19$, $p = .26$.

Discussion

The data supported our predictions: First, naive analysts made a highly reliable distinction between generic and specific listener

² Full analysis procedures are available from Janet Bavelas.

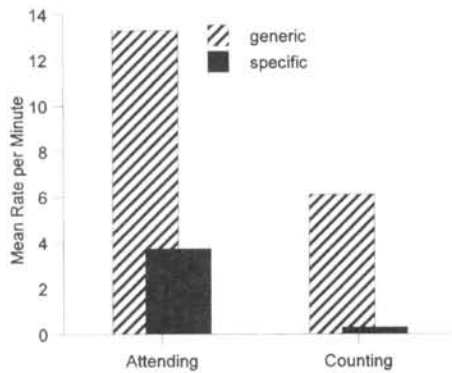


Figure 1. Effect of experimental condition on the rates of generic and specific listener responses in Experiment 1. The three attending conditions (listening, summarizing, and retelling) did not differ and are grouped together here ($n = 29$ dyads) for contrast with the counting condition ($n = 10$ dyads).

responses. Second, whereas generic responses occurred throughout the first minute, specific responses tended to occur later in this segment. Our interpretation is that specific responses follow the narrative structure of this kind of story. As we mentioned above, close-call stories typically begin with background information; listeners who have never heard the story before cannot plausibly make specific responses until they have enough information. In addition, specific responses serve an important function closer to the end of the story, where they can enhance the drama by illustrating the appropriate reaction to the events. At the climax of such stories, the listener can become a co-narrator.

Third, listeners who were distracted from the narrative made significantly fewer generic and specific responses than listeners in any of the attending conditions. It is important to note that, in the three attending conditions, cognitive load in the form of different levels of listening demands did not affect listeners' response rates. That is, even being required to memorize the details of the story for retelling did not suppress the listeners' responses, whereas a task unrelated to the story had a significant effect. Our interpretation is that as long as listeners are attending to the meaning of the story, they are capable of handling several simultaneous demands and responses, but when they are not able to attend to narrative meaning, their capacity is severely limited. Specific responses, which are more tightly connected to story content, appeared to be more affected by experimental condition than were generic responses, but this difference was not statistically significant. However, because specific responses tended to occur later, it may have been that analyzing only the first minute of the story did not provide an adequate comparison.

Finally, distracting the listener affected the overall quality of the narrator's storytelling, indicating a reciprocal effect of listener on narrator. No matter how good the story plot is, a good listener is crucial to telling it well.

Experiment 2

Our second experiment was a replication of Experiment 1, designed to confirm and extend the findings as well as to address some possible alternative explanations for these findings. First, the

counting task may have removed the listeners almost completely from the social interaction. A good strategy for the person counting would be to act as though the narrator were not there. Although it is evident from their generic responses that these listeners did not detach completely, one could argue that the overall decrease in listener responses was a function of relative disengagement from the social interaction. A second and related point can be made about the narrators, who knew that the listeners would be distracted from their storytelling. They may have told their stories differently (e.g., more poorly) because of this knowledge, that is, there was no point in trying to tell a good story. Either or both of these factors would confound our independent variable.

Third, we suspected that sometimes a listener had finished the counting task before the narrative was over (which was one reason we analyzed only the first minute). Other times, it seemed that the listener would briefly abandon the counting task and tune in to the story (because he or she started responding), but we had no independent way of confirming this. Both possibilities would mean that some listeners in the counting condition had been able to respond more than those who performed the counting task consistently throughout the narrative.

We addressed all of these concerns by devising a new counting task that prevented the listener from attending to the meaning of the narrative while keeping him or her connected to the narrator's words. The listener had to count the number of words the narrator said that began with the letter *t* and press a button each time an initial *t* was detected (see details below). This task required that the listener attend extremely closely to the narrator's words but not to their meaning, and the button-pressing gave us a manipulation check. These listeners were highly attentive, in the sense of looking constantly at the narrator and listening intently, but they were not attending to the narration. They were listening for letters in individual words, not to the meaning of the narrative. Also, in this experiment, the narrators were unaware of the listener's exact instructions, so they could not be affected by them.

There were three other changes of note: (a) We wanted to know more precisely how distracted listeners affected the narrator's story telling, so we developed an analysis of the specific narrative features that characterize good and poor story endings. (b) Because the three different attending conditions in Experiment 1 were indistinguishable from each other, there was only one control group, the middle (summarizing) condition. (c) To capture the full narrative, we analyzed listener responses for the entire story, not just the first minute.

Method

Participants

Sixty-eight participants from the University of Victoria first-year psychology classes participated in 34 dyads, for research credit. All were strangers to each other. There were equipment problems with four dyads, and six did not meet our criterion for *t*-counting (see *Procedure*, below). This left the planned total of 24 dyads for analysis, 12 randomly assigned to each condition. There were 24 women and 24 men, with the gender mix (female-female, male-male, female-male) counterbalanced across experimental condition. All participants consented to being videotaped and gave permission for analysis of their tapes afterward.

Equipment

The video recording and analysis equipment were the same as in Experiment 1. The split-screen layout of the participants was also the same, with the addition of a third camera to record the button-pressing; see below.

Procedure

Dyads were randomly assigned to listen to the narrator's story (narrative condition) or to the narrator's words (word-counting condition). One member of the pair was randomly chosen to tell his or her close-call story, and the experimenter instructed the listener while the narrator was in a different room. In both conditions, the narrator knew only that the listener was going to be listening for something in the narrative. In the word-counting (experimental) condition, the listeners were asked to count the number of words spoken by their narrator that began with the letter *t*. To monitor the listeners' performance, we asked them to press a button held in their laps, under the table, every time they heard a word that began with a *t*. The button was connected to a small light at the end of the table, which was not visible to either participant but was captured on a third camera and split onto the upper right corner of the videotape. With this real-time record, we could check each listener's *t* count against our own. We subsequently analyzed only dyads where the listeners caught at least 60% of the words beginning with *t*. (The main reason for the relatively low cut-off level was that most listeners missed words that began with *th*.)

The narrative condition was essentially the same as the summarizing condition from Experiment 1. The listener was to "listen to the story so that, if you had to, you could summarize the gist or main point of it to someone else."

At the end of the experiment, the narrators learned exactly what their listeners had been doing. As part of the debriefing for the word-counting condition, narrators had a chance to try out the role of listener so that they would understand the nature of the task and its effects on the listener's behavior.

Analyses

Listener responses. The entire narrative for all 24 dyads was analyzed for generic and specific responses by the method described in Experiment 1. Both analysts were unaware of hypotheses and conditions, and the video monitor they used had a fixed cover over the flashing light. Their overall reliability was 99.6% agreement.

Story endings. We analyzed the story endings (what happened right after the climax of the close call) on four dichotomous scales:

1. Was the pace of the ending *appropriate* or *abrupt*?

Some narrators chose to end the story immediately after the climax, without further detail. In this case, it would be appropriate to do so at the same pace as the story was told. The pace of the ending was rated as abrupt when the pace changed, either speeding up, as if to get it over, or fading away.

2. Did the narrator *appropriately extend* the denouement or just *talk on and on*?

If he or she did not end the story once the climax had occurred, the narrator might continue with appropriate detail that elaborated the story or "milked" the drama of the story. Alternatively, he or she might just continue talking as if looking for a way to end the story, often repeating information.

3. Was the ending *choppy* or *not*?

After the climax of the story, the narrator might continue at an uneven pace, pausing or becoming disfluent during or between sentences in such a way as to produce noticeable gaps and a choppy pace.

4. Did the narrator try to *justify* or *explain* the closeness of the close call or *not*?

Justifying was present when the narrator attempted to emphasize the

obvious risk or danger, as if trying to convince the listener that it really was a close call.

The score for each story ending was the total number of negative features (i.e., abrupt, talks on and on, choppy, justifying). The total number of negative story-ending features could range from 0 to 3, because two of the four features were mutually exclusive (the narrator could not both end the story abruptly and also talk on and on). Two stories in the narrative condition were not analyzed because of poor audio quality.

Two analysts worked independently and then collaborated to produce one set of ratings, which were compared to a third independent analyst's ratings for reliability ($r = .69$). The same two analysts also rated the story endings from Experiment 1; agreement with the third analyst for a sample of 20 of these stories was $r = .76$.

Results

Serial Position of Specific and Generic Responses

Using the 12 dyads in which both generic and specific responses occurred, we analyzed the average serial position of the two kinds of responses for the entire narrative. Specific responses occurred significantly later in the narratives ($M = 18.46$, $SD = 12.11$) than did generic responses ($M = 14.01$, $SD = 8.60$); $t(11) = 2.85$, $p = .016$. This extends the findings of Experiment 1, for which we had analyzed only the first minute.

Effect of Condition on Rate of Generic and Specific Responses per Minute

Listeners in the word-counting condition made significantly fewer specific responses per minute ($M = .08$, $SD = .29$) than did participants in the narrative condition ($M = 2.21$, $SD = 1.21$); $t(22) = 5.94$, $p < .0001$ (see Figure 2).

Parametric testing did not detect a significant effect of experimental condition on generic responses, narrative: $M = 8.91$, $SD = 2.67$; word-counting: $M = 7.44$, $SD = 5.95$; $t(22) = 0.78$, $p = .45$. However, as the standard deviations suggest, this was due to a few outliers in the word-counting condition who managed to make a high number of generic responses. The appropriate non-parametric test revealed that generic responses were in fact signif-

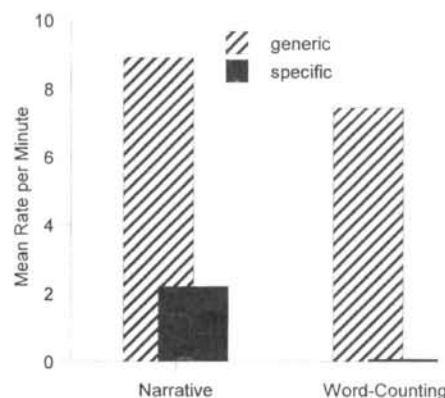


Figure 2. Effect of experimental condition on the rates of generic and specific listener responses in Experiment 2. There were 12 dyads in each condition.

icantly reduced by the word-counting task (Mann-Whitney $U = 41$; $p = .034$).

As can be seen, the word-counting task reduced the rate of generic responses to 80% of their rate in the narrative condition, whereas specific responses dropped to less than 5% of the rate in the narrative condition. This experimental condition, which was intended to be more demanding than that of Experiment 1, virtually eliminated specific responses. There was a substantial difference in effect sizes; for specific responses $\eta^2 = .62$, for generic responses $\eta^2 = .03$. (There was no significant interaction between experimental condition and generic versus specific response, treated as a second factor, but this is probably because of the heterogeneity of variance, noted above, and the small sample size, which did not provide enough power to detect an effect.)

Effect of Experimental Condition on Quality of Narration

The stories told in the word-counting condition had more negative features ($M = 2.25$; $SD = 1.22$) than those in the narrative condition ($M = .40$, $SD = .70$); $t(20) = 4.25$, $p < .0001$. We used the ratings of the more experienced pair of analysts, but because of the moderate reliability, we repeated the analysis using only the third analyst's ratings; these were also significant. We also replicated the effect of the experimental condition on negative features of story endings in the Experiment 1 data, counting: $M = 1.75$, $SD = .89$; attending: $M = .69$, $SD = .89$; $t(35) = 2.98$, $p = .005$. The effect of condition on the four different negative features in both experiments is shown in Figure 3.

We can illustrate these effects with the end of a story told in the word-counting condition by a particularly skillful story-teller. He began his story by explaining that he was working as a logger one summer when he and his partner misjudged the height of a tree that was going to fall into the narrow open corridor in which they were working. There was no place to escape on either side, so as they realized their danger, they could only try to outrun the falling tree. The rest of his story is transcribed below. During the ending, as throughout the story, the listener nodded and occasionally smiled but made no specific responses (i.e., no facial displays of fear or concern, no verbal interjections, etc.). The story reached its dra-

matic climax at the end of the third sentence below, but the only response was a slight nod and smile. Notice the change in the narration after that point:

So this tree's falling, falling, falling. And he was ahead of me, and I was behind him, and *just* the end of the tree clipped my foot. And it felt like, like a *whip* hitting my foot. And so ah after I, I mean, I saw it fall and we both go *diving* into the thing cause we knew—I mean, I don't know how exciting that is but afterwards, ah, I mean, we chuckled about it at lunch. Cause it's always funny if you don't get landed on, sure it was a hoot, but (stylized laugh). Um. I just thought that was, ah, that was funny that, ah. Like *usually*, the easy way to go out is go to either side, and that way it'll fall and you're on either side. But since we had no escape route, we knew it was comin' at us, so we had to run for our lives basically, which puts a little excitement into the job too, cause it's fun, rappelling down trees and stuff and, and what-not. So . . . that's all!

Immediately after the climax ("And it felt like, like a *whip* hitting my foot"), his narration illustrates three negative features of story endings: The rest of the excerpt illustrates talking on and on, because none of it adds to the excitement of the story (e.g., needlessly reiterating the problem of no escape route, as if trying get the point across, and adding irrelevant information about the job in general, such as "rappelling down trees"). His previously smooth and skillful delivery became choppy and uneven in pace, with broken phrases and unfinished sentences. He slowed down for disfluencies and filler words ("so ah," "I mean") and speeded up again for new explanations. Finally, he justified the story as a close call by pointing out the obvious danger ("we had to run for our lives basically") while at the same time seeming almost to apologize for or retract the story ("I don't know how exciting that is").

Other Variables: Gender and Story Length

As in Experiment 1, there was no significant effect of listener gender on the rate of generic, specific, or total listening responses; $t(22) = 0.44$, 0.64 , and 0.23 , all $ps > .52$. The mean story length was 2 min 28 s ($SD = 1$ min 35 s) and was not affected by condition, $t(22) = 1.42$, $p = .18$.

Discussion

The results of Experiment 2 replicated and refined the results of Experiment 1. The mean serial position of specific responses was later than that of generic responses for the entire story, not just the first minute. The experimental manipulation eliminated some possible confounds and expressed our theory more precisely: The listener must be not merely attending but attending at the level of the narrative. The more demanding experimental condition revealed that listeners attending to individual words could not or did not make as many specific responses as participants who were attending to the overall narrative. Generic responses were also affected, although to a considerably lesser degree. This difference is consistent with our definition of specific responses as tightly connected to narrative content. They must occur at the right time, and they must be shaped for that moment in the story. Finally, a microanalysis of features that make the climax of a close-call story well or poorly told revealed that the stories faltered or fell flat when they were told to listeners who were attending closely to the individual words but not to the narrative itself.

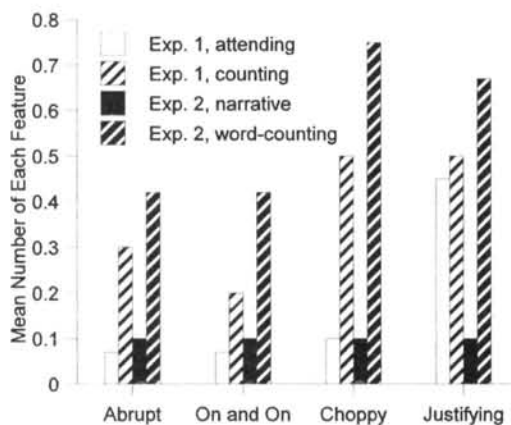


Figure 3. Effect of experimental condition on four negative features of story endings in Experiments 1 and 2. (See Method, Experiment 2, for explanation of the features.)

General Discussion

These two experiments confirmed all of our predictions and demonstrated important differences between generic and specific listener responses as well as the importance of the listener to the narrator. There were four major findings:

First, even though distinguishing between the two kinds of responses requires interpretation, independent analysts who were unaware of our purpose and hypotheses made the distinction with very high reliability. We have had similarly high agreement for other meaning-based analyses (e.g., Bavelas et al., 1986, 1995; Chovil, 1991/1992; Coates, 1991), which should be encouraging to other researchers who want to move in this direction while maintaining experimental standards.

Second, as we and Goodwin (1986) had predicted, specific responses occurred significantly later in the story than did generic responses, which occurred throughout the narrative. Listeners do not or cannot make specific responses until they have enough information about the narrative. Also, specific responses are especially appropriate to exciting or risky events, which came later in the narratives studied here. Thus, specific responses are much more sensitive to the intricacies of the narrative than are generic responses. Their different relationships to the narrative itself is our first piece of evidence that listener responses should not be treated as a homogeneous group.

Third, listeners who were experimentally distracted from the narrative still made some generic responses but almost no specific responses. The listener had to be following the meaning of the story closely to be able to insert precise and specific responses at appropriate points in the narrative. The lack of difference among the three attending conditions in Experiment 1 showed that increasing the level of cognitive demand on the listener did not matter as long as he or she was still attending to meaning. A close connection to the narrative rather than cognitive demand was the crucial factor for making listener responses. The findings of Experiment 2 confirmed more precisely that the listener has to be listening not only to the speaker's words but also to the speaker's meaning.

It is important not to interpret our findings as showing that specific responses are in some sense preferable to generic responses or that specific responses are a better way of listening. These are simply two different kinds of responses, with different functions in conversation. Our data showed that generic responses are the most common kind of listener response in normal listening conditions and that they occurred throughout the narrative. Nor should the fact that distracted listeners in our unusual experimental conditions could still make some generic responses suggest to the reader that generic responses are normally evidence of inattentiveness. Indeed, the generic responses that occurred in our counting conditions were well-timed and appropriate. (We are currently analyzing precisely what elicits these responses.)

Fourth, the same experimental conditions that distracted the listener from the narrative also strongly affected the narration itself. Narrators who told close-call stories to distracted listeners or to listeners who were attending closely to the speaker's words rather than their narrative meaning told them less well overall and particularly poorly at what should have been the dramatic conclusion. Their story endings were abrupt or choppy, or they circled around and retold the ending more than once, and they often

justified their story by explaining the obvious close call. One highly plausible reason is that the relative absence of listener responses, particularly specific responses, caused the narration to falter. That is, in our model, part of the narrators' problem was that the listeners were not making their contribution to the narrative. Equally important is the fact that, in contrast to normal listeners, they were not "in sync" with their narrator, helping them moment-by-moment to finish the story smoothly and effectively. In the autonomous view, their close participation would be irrelevant because only the narrator's story plot and skill would matter. We will return to some practical implications of this effect below.

It would be desirable to test this causal hypothesis (that it was the absence of specific responses that made the stories falter) with a mediation analysis (Baron & Kenny, 1986; Kenny, Kashy, & Bolger, 1998), but that is not possible, for an interesting reason. The available models assume a linear causal sequence in which the experimental condition affected the listener's behavior, which in turn affected the narrator's behavior. Our collaborative model assumes a constant *reciprocal* influence between narrators and listeners; we assume that ordinarily the narrator and listener work together moment by moment to produce a good story. The presence or absence of appropriate listener responses would affect the quality of narration, but the quality of narration would also affect the quality of listener responses. That is, a responsive listener would improve the narrative, which would increase the likelihood that the listener would continue to respond, and so forth. An unresponsive listener would dampen the narrative, which would make the narrative less likely to elicit responses, and so forth. To the extent that this reciprocal process is happening, then the data would violate several assumptions of mediation analysis: reverse causal effects (Baron & Kenny, 1986, p. 1177), proximal mediation (Kenny et al., 1998, p. 261), and multicollinearity (Kenny et al., 1998, p. 263). This was true in our data, and we know of no alternative analysis that is suited to a micro-reciprocal causal model. Therefore, however plausible, our causal interpretation cannot be statistically supported at present.

To our knowledge, these are the first experimental studies of what listeners do in spontaneous narratives told in face-to-face dialogue. The task is similar to the kind of dyadic story-telling that occurs in many natural settings (e.g., among friends or at parties). Because the listener could not know anything about the story in advance, the dialogues were, as we intended, quite asymmetrical. In any monologic or autonomous view of conversation, "mere listeners" should have nothing to contribute. Yet these listeners who had no formal speaking role and no prior knowledge of the story became co-narrators, in two senses: First, through their specific responses, they contributed vivid and helpful components to the narrator's story, which the narrator might otherwise have made him- or herself. Second, it is apparent from the deleterious effect of a distracted listener on the quality of the narration that listeners assist the narrator in telling a good story. Listeners may not be equal co-narrators, but they are essential.

Pasupathi, Stallworth, and Murdoch (1998) adapted our procedure to study the effects of a distracted listener on the narrators' memory of a story that had been supplied to them by the experimenter. When a confederate listener counted the number of words that began with *th*, the narrators actually remembered less of the material they had told. They also rated the experience as less pleasant, which corresponds to our observations, particularly those

from Experiment 2. Tatar (1997) used a similar distraction task (with participants rather than confederates as listeners). She found that speakers talking about a life experience that had produced pride or shame rated themselves as less engaged and told shorter stories when the listener was distracted by counting *ths*.

Our findings have unexpected implications for many kinds of interviewing. For example, we have proposed (Routledge, Bavelas, McGee, & Wade, 1994) that our distracted listeners may resemble some psychotherapists when they assume that they are "mere listeners" to their clients' narratives and are preoccupied with a focus that is other than the narrative itself (e.g., with body language, diagnosis, or self-esteem issues). Similarly, the interviewer (listener) in job interviews, police interviews, and survey research often tries to be unresponsive in an effort not to influence the interviewee. Our results suggest that their lack of response could have a strong (and negative) influence. Collaborative effects may be inevitable in face-to-face dialogue.

Altogether, these results strongly support our view that face-to-face dialogue is distinct from monologue or written text because it includes microsocial processes of coordination and collaboration in addition to the already well-studied processes of language production and comprehension. Even in highly asymmetrical dialogues, speaker and listener roles are not fixed and separate. Rather, their relationship is reciprocal and collaborative, in that the narrator elicits responses from the listener and the listener's responses affect the narrator. In spontaneous storytelling, the interlocutors interact together to produce the narrative. We agree with those (e.g., Clark, 1994, 1996; Goodwin, 1979; Goodwin & Goodwin, 1987) who propose that narrative is a joint activity and does not belong to either speaker or listener alone. Dialogue is not simply information transmission between individuals but is a reciprocal process of co-construction. The essential contribution of listeners must be included to understand language use in face-to-face dialogue.

References

- Atkinson, J. M., & Heritage, J. (1984). *Structures of social action: Studies in conversational analysis*. Cambridge, England: Cambridge University Press.
- Baron, R. M., & Kenny, D. A. (1986). The moderator-mediator variable distinction in social psychological research: Conceptual, strategic, and statistical considerations. *Journal of Personality and Social Psychology*, 51, 1173-1182.
- Bavelas, J. B. (1990). Nonverbal and social aspects of discourse in face-to-face interaction. *Text*, 10, 5-8.
- Bavelas, J. B. (1994). Gestures as part of speech: Methodological implications. *Research on Language and Social Interaction*, 27, 201-221.
- Bavelas, J. B., Black, A., Chovil, N., & Mullett, J. (1990). *Equivocal communication*. Newbury Park, CA: Sage.
- Bavelas, J. B., Black, A., Lemery, C. R., & Mullett, J. (1986). "I show how you feel." Motor mimicry as a communicative act. *Journal of Personality and Social Psychology*, 50, 322-329.
- Bavelas, J. B., & Chovil, N. (1997). Faces in dialogue. In J. Russell & J.-M. Fernandez-Dols (Eds.), *The psychology of facial expression* (pp. 334-346). Cambridge, England: Cambridge University Press.
- Bavelas, J. B., & Chovil, N. (2000). Visible acts of meaning: An integrated message model of language in face-to-face dialogue. *Journal of Language and Social Psychology*, 19, 163-194.
- Bavelas, J. B., Chovil, N., Coates, L., & Roe, L. (1995). Gestures specialized for dialogue. *Personality and Social Psychology Bulletin*, 21, 394-405.
- Bavelas, J. B., Hutchinson, S., Kenwood, C., & Matheson, D. H. (1997). Using face-to-face communication as a standard for other communication systems. *Canadian Journal of Communication*, 22, 5-24.
- Bruner, J. (1990). *Acts of meaning*. Cambridge, MA: Harvard University Press.
- Chovil, N. (1991/1992). Discourse-oriented facial displays in conversation. *Research on Language and Social Interaction*, 25, 163-194.
- Clark, H. H. (1994). Discourse in production. In M. A. Gernsbacher (Ed.), *Handbook of psycholinguistics* (pp. 985-1021). San Diego: Academic Press.
- Clark, H. H. (1996). *Using language*. Cambridge, England: Cambridge University Press.
- Clark, H. H., & Wilkes-Gibbs, D. (1986). Referring as a collaborative process. *Cognition*, 22, 1-39.
- Coates, L. (1991). *A collaborative theory of inversion: Irony in dialogue*. Unpublished master's thesis, Department of Psychology, University of Victoria, British Columbia, Canada.
- Coates, L., & Johnson, T. (in press). Toward a social theory of gender. In H. Giles & P. Robinson (Eds.), *Handbook of language and social psychology* (2nd ed.). New York: Wiley.
- Engle, R. A., & Clark, H. H. (1995, March). *Using composites of speech, gestures, diagrams and demonstrations in explanations of mechanical devices*. Paper presented at the American Association for Applied Linguistics Conference, Long Beach, CA.
- Feldman, C. F. (1971). The interaction of sentence characteristics and mode of presentation in recall. *Language and Speech*, 14, 18-25.
- Fillmore, C. (1981). Pragmatics and the description of discourse. In P. Cole (Ed.), *Radical pragmatics* (pp. 143-166). New York: Academic Press.
- Fivush, R. (1991). The social construction of personal narratives. *Merrill-Palmer Quarterly*, 37, 59-82.
- Flexner, S. B. (Ed.). (1983). *Random House Unabridged Dictionary* (2nd ed.). New York: Random House.
- Goodwin, C. (1979). The interactive construction of a sentence in natural conversation. In G. Psathas (Ed.), *Everyday language: Studies in ethnomethodology*. New York: Irvington Publishers.
- Goodwin, C. (1981). *Conversational organization: Interaction between speakers and hearers*. New York: Academic Press.
- Goodwin, C. (1986). Between and within: Alternative sequential treatments of continuers and assessments. *Human Studies*, 9, 205-217.
- Goodwin, C., & Goodwin, M. H. (1987). Concurrent operations on talk: Notes on the interactive organization of assessments. *International Pragmatics Association Papers in Pragmatics*, 1, 1-54.
- Kenny, D. A., Kashy, D. A., & Bolger, N. (1998). Data analysis in social psychology. In D. T. Gilbert, S. T. Fiske, & G. Lindzey (Eds.), *The handbook of social psychology* (4th ed., pp. 233-265). Boston: McGraw-Hill.
- Kintsch, W., & van Dijk, T. A. (1978). Toward a model of text comprehension and production. *Psychological Review*, 85, 363-394.
- Krauss, R. M., Garlock, C. M., Bricker, P. D., & McMahon, L. E. (1977). The role of audible and visible back-channel responses in interpersonal communication. *Journal of Personality and Social Psychology*, 35, 523-529.
- Krauss, R. M., & Weinheimer, S. (1966). Concurrent feedback, confirmation, and the encoding of referents in verbal communication. *Journal of Personality and Social Psychology*, 4, 343-346.
- Kraut, R. E., Lewis, S. H., & Swezey, L. W. (1982). Listener responsiveness and the coordination of conversation. *Journal of Personality and Social Psychology*, 43, 718-731.
- Labov, W. (1972). The transformation of experience in narrative syntax. In W. Labov (Ed.), *Language in the inner city*. Philadelphia: University of Pennsylvania Press.
- Leavitt, H. J., & Mueller, R. A. (1951). Some effects of feedback on communication. *Human Relations*, 4, 401-410.

- Linell, P. (1982). *The written language bias in linguistics*. Linköping, Sweden: University of Linköping.
- Mandler, J. M. (1987). A code in the node: The use of a story schema in retrieval. *Discourse Processes*, 1, 14-35.
- Marche, T. A., & Peterson, C. (1993). On the gender differential use of listener responsiveness. *Sex Roles*, 29, 795-816.
- Miller, G. A. (1951). *Language and communication*. New York: McGraw-Hill.
- Neisser, U. (1982). Snapshots or benchmarks? In U. Neisser (Ed.), *Memory observed* (pp. 43-48). San Francisco: Freeman.
- Pasupathi, M., Stallworth, L. M., & Murdoch, K. (1998). How what we tell becomes what we know: Listener effects on speakers' long-term memory for events. *Discourse Processes*, 26, 1-25.
- Phillips, B., & Bavelas, J. B. (2000, June). *Comparing collaborative versus autonomous models of conversation in computer-mediated communication*. Paper presented at the meeting of the International Communication Association, Acapulco.
- Polanyi, L. (1989). *Telling the American story*. Cambridge, MA: MIT Press.
- Routledge, R., Bavelas, J. B., McGee, D., & Wade, A. (1994, July). Therapy and research. Workshop presented at Australian Family Therapy Association, Sydney, Australia.
- Rumelhart, D. (1975). Notes on a schema for stories. In D. Bobrow & A. Collins (Eds.), *Representation and understanding* (pp. 237-272). New York: Academic Press.
- Sacks, H. (1974). An analysis of the course of a joke's telling in conversation. In R. Bauman & J. Scherzer (Eds.), *Explorations in the ethnog-*

- raphy of speaking* (pp. 337-353). Cambridge: Cambridge University Press.
- Sacks, H., Schegloff, E. A., & Jefferson, G. (1974). A simplest systematics for the organization of turn-taking in conversation. *Language*, 50, 696-735.
- Schober, M., & Clark, H. H. (1989). Understanding by addressees and overhearers. *Cognitive Psychology*, 21, 211-232.
- Shannon, C. E., & Weaver, W. (1949). *A mathematical theory of communication*. Urbana: University of Illinois Press.
- Streeck, J. (1994). Gesture as communication II: The audience as co-author. *Research on Language and Social Interaction*, 27, 239-267.
- Tatar, D. (1997). *Social and personal consequences of a preoccupied listener*. Unpublished doctoral dissertation, Department of Psychology, Stanford University, Stanford, CA.
- Watzlawick, P., Beavin Bavelas, J., & Jackson, D. D. (1967). *Pragmatics of human communication*. New York: Norton.
- Williams, E. (1977). Experimental comparisons of face-to-face and mediated communication: A review. *Psychological Bulletin*, 84, 963-976.
- Yngve, V. H. (1970). On getting a word in edgewise. *Papers from the Sixth Regional Meeting of the Chicago Linguistic Society* (pp. 567-578). Chicago: Chicago Linguistic Society.

Received December 18, 1998

Revision received March 29, 2000

Accepted March 29, 2000 ■

United States Periodicals Service
Statement of Ownership, Management, and Circulation

1. Publication Title: *Journal of Personality and Social Psychology*

2. Publication Number: 0022-3816

3. Filing Date: October 2000

4. Issue Frequency: Monthly

5. Number of Issues Published Annually: 12

6. Annual Subscription Price: \$170, Indiv \$340, Inst \$512

7. Complete Mailing Address of Known Office of Publication (Not printer) (Street, city, county, state, and ZIP+4):
750 First Street, NE - Washington, DC 20002-4242

8. Complete Mailing Address of Headquarters or General Business Office of Publisher (Not printer):
750 First Street, NE - Washington, DC 20002-4242

9. Full Names and Complete Mailing Addresses of Publisher, Editor, and Managing Editor (Not for leave blank):
Publisher (Name and complete mailing address):
American Psychological Association
750 First Street, NE
Washington, DC 20002-4242
Editor (Name and complete mailing address):
Chester A. Insko, PhD, Dept of Psychology, Univ of NC, Chapel Hill, NC 27599
Arie Kruglanski, PhD, Dept of Psychology, Univ of MD, College Park, MD 20742
Ed Streiner, PhD, Dept of Psychology, Univ of IL, 603 East Daniel, Champaign, IL 61820
Managing Editor (Name and complete mailing address):
Susan Knapp
American Psychological Association
750 First Street, NE - Washington, DC 20002-4242

10. Owner (Do not leave blank. If the publication is owned by a corporation, give the name and address of the corporation immediately followed by the names and addresses of all stockholders owning or holding 1 percent or more of the total amount of stock. If not owned by a corporation, give the names and addresses of the individual owners. If owned by a partnership or other unincorporated firm, give its name and address as well as those of each individual owner. If the publication is published by a foreign corporation, give its name and address.)

Full Name: American Psychological Association
Complete Mailing Address: 750 First Street, NE
Washington, DC 20002-4242

11. Known Bondholders, Mortgagees, and Other Security Holders Owning or Holding 1 Percent or More of Total Amount of Bonds, Mortgages, or Other Securities. If none, check box: ☒ None

Full Name: None
Complete Mailing Address:

12. Tax Status (For completion by nonprofit organizations authorized to mail at nonprofit rates) (Check one):
☐ The purpose, function, and nonprofit status of this organization and the exempt status for federal income tax purposes:
☐ Has Not Changed During Preceding 12 Months
☒ Has Changed During Preceding 12 Months (Publisher must submit explanation of change with this statement)

13. Publication Title: *Journal of Personality and Social Psychology*

14. Issue Date for Circulation Data Below: August 2000

15. Extent and Nature of Circulation

Average No. Copies Each Issue During Preceding 12 Months		No. Copies of Single Issue Published Nearest to Filing Date
a. Total Number of Copies (Net press run)		
(1) Paid and/or Requested Circulation	6,377	6,397
(2) Paid Outside-County Subscriptions Stated on Form 3541 (Include advertiser's proof and exchange copies)	1,186	1,168
(3) Paid In-County Subscriptions Stated on Form 3541 (Include advertiser's proof and exchange copies)	4,055	3,630
(4) Sales Through Dealers and Carriers, Street Vendors, Counter Sales, and Other Non-USPS Paid Distribution		
(5) Other Classes Mailed Through the USPS		
b. Total Paid and/or Requested Circulation (Sum of 15b(1), (2), (3), (4), and (5))		
	5,241	4,826
c. Free Distribution by Mail (Carriers or other means)		
(1) Outside-County as Stated on Form 3541		
(2) In-County as Stated on Form 3541	406	329
(3) Other Classes Mailed Through the USPS		
d. Free Distribution Outside the Mail (Carriers or other means)		
(1) Total Free Distribution (Sum of 15d(1) and 15d(2))	406	329
e. Total Distribution (Sum of 15b and 15d)		
	5,647	5,155
f. Copies not Distributed		
	730	1,242
g. Total (Sum of 15e and 15f)		
	6,377	6,397
h. Paid and/or Requested Circulation (15b, divided by 12, times 100)		
	92.8	93.6

16. Publication of Statement of Ownership:
☒ Publication required. Will be printed in the December 2000 issue of this publication.
☐ Publication not required.

17. Signature and Title of Editor, Publisher, Business Manager, or Owner:
Ed Streiner, PhD, Dept of Psychology, Univ of IL, 603 East Daniel, Champaign, IL 61820
Date: 8/8/00

I certify that all information furnished on this form is true and complete. I understand that anyone who furnishes false or misleading information on this form or who omits material or information requested on the form may be subject to criminal sanctions (including fines and imprisonment) and/or civil sanctions (including civil penalties).

Instructions to Publishers

- Complete and file one copy of this form with your postmaster annually on or before October 1. Keep a copy of the completed form for your records.
- In cases where the stockholder or security holder is a trustee, include in items 10 and 11 the name of the person or corporation for whom the trustee is acting. Also include the names and addresses of individuals who are stockholders who own or hold 1 percent or more of the total amount of bonds, mortgages, or other securities of the publishing corporation. In item 11, if none, check the box. Use blank sheets if more space is required.
- Be sure to furnish all circulation information called for in item 15. Free circulation must be shown in items 15c, e, and f.
- Item 15b, Copies not Distributed, must include (1) newspaper copies originally stated on Form 3541, and returned to the publisher; (2) estimated returns from news agents; and (3) copies for office use, reference, and all other copies not distributed.
- If the publication had Periodicals authorization as a general or requester publication, this Statement of Ownership, Management, and Circulation must be published. It must be printed in any issue in October or, if the publication is not published during October, the first issue printed after October.
- In item 16, indicate the date of the issue in which this Statement of Ownership will be published.
- Item 17 must be signed.

Failure to file or publish a statement of ownership may lead to suspension of Periodicals authorization.

PS Form 3526, October 1999 (Rev 09/99)