

# Spontaneous Trait Inferences Are Bound to Actors' Faces: Evidence From a False Recognition Paradigm

Alexander Todorov and James S. Uleman  
New York University

A false recognition paradigm showed that spontaneous trait inferences (STIs) are bound to the person performing a trait-implying behavior. In 6 experiments, participants memorized faces and behavioral sentences. When faces were paired with implied traits in a recognition test, participants falsely recognized these traits more often than unrelated traits paired with the same faces or the same traits paired with familiar faces. The effect was obtained for a large set of behaviors (120), each presented for 5 s, and for behaviors that participants did not subsequently recognize or recall. Antonyms of the implied traits were falsely recognized less often than unrelated traits, suggesting that STIs have extended implications. Explicit person–trait judgments predicted both false recognition and response times for implied traits.

Social life is replete with opportunities to form impressions of strangers. Occasionally people do this intentionally—at job interviews, meetings with new colleagues, and social mixers. More often people do this spontaneously, without having a particular goal or even a general impression-formation intention in mind and without becoming aware that they have made an inference. By now, the occurrence of spontaneous trait inferences (STIs) is fairly well established (see Uleman, Newman, & Moskowitz, 1996). But at least two important issues remain unresolved.

The first concerns what these STIs refer to. When most European Americans<sup>1</sup> read that “Anna solved the mystery half way through the book,” the trait concept *clever* is activated. Does this concept refer to Anna, to the behavior of solving the mystery early, or to both the actor and the behavior? When trait inferences are unintentional and even unconscious, how can one know what they refer to? The second issue concerns the implications of STIs. Having inferred spontaneously that Anna and/or her behavior is clever, does one also know that she is not dull? How far beyond the information given does one go?

In addition, we were interested in developing a better measure of STIs—one that would reveal STIs occurring in the face of explicit task demands, one that would reveal STIs in long term memory (LTM), and one that would yield large effect sizes. The false recognition procedure used in the six studies reported here seems to have these advantages.

## The False Recognition Paradigm

If an inference is made unintentionally and unconsciously, in what sense can it refer to anything? One way is for it to be linked

or bound to the actor in LTM. Links in memory between an actor and a trait implied by the actor's behavior can be revealed either by explicit or by implicit memory tasks. Such links seem to be clearly supported by prior research using implicit tasks (Carlston & Skowronski, 1994; Carlston, Skowronski, & Sparks, 1995). However, the evidence on actor–trait links from explicit memory tasks is weak (Uleman et al., 1996).

The only explicit paradigm used to specifically study actor–trait links in STIs has been the cued recall paradigm. Participants study a series of trait-implying sentences for a subsequent memory test. Then, after a filler task, explicit memory for each sentence is cued with one of several cues. Actor–trait links are demonstrated when trait cues facilitate recall of actors, with recall of behaviors controlled for. In the typical STI study, such as Winter and Uleman (1984), in which participants tried to memorize one sentence and behavior per actor, trait cues facilitated sentence recall. But when recall of behaviors was taken into account, there was no consistent evidence of actor–trait links in any of these studies (see Uleman et al., 1996, p. 246).

In contrast to the cued recall paradigm, which relies exclusively on trait-to-actor links in LTM, the false recognition paradigm takes advantage of both trait-to-actor and actor-to-trait links. If STIs really represent inferences about actors, then the information should be organized more by actors than by traits. The cued recall paradigm (in which traits are cues) does not take advantage of this organization.

In the first (study) phase of the false recognition paradigm, participants are presented with persons' faces and behavioral sentences (one sentence per face).<sup>2</sup> In the study phase, participants are

---

Alexander Todorov and James S. Uleman, Department of Psychology, New York University.

We thank Ran Hassin and John Skowronski for their comments on drafts of this article and Emily Thaden, Kevin Young, and Lauren Pagano for their help in conducting the experiments.

Correspondence concerning this article should be addressed to Alexander Todorov, who is now at the Department of Psychology, Green Hall, Princeton University, Princeton, New Jersey 08544-1010. E-mail: atodorov@princeton.edu

---

<sup>1</sup> Zárate, Uleman, and Voils (2001) showed that Latinos are less likely to make STIs than are Anglos. There is also evidence that STIs are more likely among individualists than collectivists (Duff & Newman, 1997).

<sup>2</sup> D'Agostino (1991) used a recognition paradigm to study STIs, but he did not present faces or use faces as retrieval prompts, did not vary actor–trait pairing at recognition, and did not include the other control conditions that we included. His questions were different from ours, and his interpretation of his results differs from ours (see Uleman et al., 1996, p. 226).

simply asked to memorize the information for a memory test later in the experiment. In the second (test) phase, participants are presented with a recognition test in which each trial consists of a face–trait pair. The participant’s task is to decide whether the trait word was presented in the behavioral sentence that described the person.

A major assumption of the current studies is that false recognition of implied traits is driven by encoding processes. The trait is inferred during the presentation of the behavior, and both the trait inference and the actual behavior are encoded as part of the actor representation. In the recognition task, the familiarity of the trait can be misattributed to a prior presentation. In other words, participants fail to discriminate between the initial inference and the information actually presented. This source-monitoring failure (M. K. Johnson, Hashtroudi, & Lindsay, 1993) should occur even if the actor’s behavior is not recalled. In the latter case, participants should falsely recognize implied traits because of their familiarity. In fact, certain conditions, such as short presentation of the behavior and a large set of behaviors, could interfere more with the encoding of the actual behavior than with the encoding of the trait inference.

If participants associate the spontaneously inferred trait with the person who performed the behavior, two effects should be observed. First, false recognition of traits should be higher when the trait implied by the behavior is paired with the actor’s face than when the trait is paired with some other face or the face is paired with some other trait. Second, participants should be slower to correctly reject an implied trait as previously seen than to correctly reject a different trait.

### Evidence for Actor–Trait Links From Implicit Tasks

Carlston, Skowronski, and their colleagues (Carlston & Skowronski, 1994; Carlston et al., 1995) have provided the most extensive evidence for actor–trait links using an implicit paradigm. In their *savings in relearning* paradigm, participants initially familiarize themselves with a series of facial photos paired with trait-implicating paragraphs. Then, after an interval as long as 1 week, they are asked to learn pairs of faces and traits. These pairs include photos presented earlier, paired with the traits implied by the behaviors that accompanied them. Such old photo–implied trait pairs are easier to learn than are new photo–implied trait pairs or old photos paired with other implied traits. That is, they show savings in relearning, because the pairings of photos and traits were evidently learned during familiarization, even though there was no instruction to infer traits or form impressions. Familiarization produces STIs, and their links to the actor photos in LTM are the basis for the savings. These effects occur not only without any explicit reference to the initial familiarization experience but also without any explicit memory for the behaviors.

The savings paradigm is the first paradigm that has provided consistent evidence for actor-linked STIs. However, there are two ambiguities in the paradigm. First, each face is paired with rich person descriptions consisting of multiple sentences. A rich description with a consistent trait theme seems to suggest making trait inferences, whatever participants’ ostensible processing goals are. Work by Park (1989) suggests that pairing multiple trait-implicating sentences with a face is sufficient to trigger impression formation goals. In fact, several experiments with the savings

paradigm failed to find differences between memory and impression formation instructions (Carlston & Skowronski, 1994). Second, in the savings paradigm, the trait inferences and the learning task demands act in the same direction. That is, trait inference facilitates the relearning task performance. More trait inference leads to better relearning. Moreover, even if a participant is aware of making a trait inference, this knowledge facilitates the relearning of face–trait pairs.

In contrast, the false recognition paradigm uses single behavioral sentences, and the trait inference and the task demands act in different directions. In this paradigm, STI leads to an incorrect response. Thus, optimal task performance (accurate recognition) and making STIs oppose each other. Further, if a participant is aware of making an inference, he or she can use this knowledge to avoid a false recognition error. Given the differences between the savings and the false recognition paradigms, the case for actor-linked STIs would be strengthened if we could provide convergent evidence in the false recognition paradigm.

### The Implications of STIs

A great deal is known about the implications that people can draw intentionally from a few traits about a person (e.g., Reeder & Brewer, 1979; Schneider, 1973). And something is known about the way that traits (and other adjectives) are organized in the lexicon (Gross, Fischer, & Miller, 1989). But very little is known about the implications that people may draw from STIs about a person. Once one spontaneously infers that a person is clever, does one infer anything else about her?

The only studies to date to look at this question have concluded that one does not do so. In three experiments, Skowronski, Carlston, Mae, and Crawford (1998; see also Mae, Carlston, & Skowronski, 1999) showed that spontaneously inferred traits did not affect ratings on denotatively unrelated traits. These studies suggest that spontaneous trait inferences do not have implications for other traits with similar or with opposite valences, as long as these traits are denotatively unrelated.

However, traits are adjectives, and Gross et al. (1989) have demonstrated that adjectives are organized in the lexicon as opposites. Thus, STIs might have implications for antonym trait ratings of an actor. Moreover, if the implied trait is encoded as part of the actor representation, then it might have implications for semantically related traits. Some of the studies below were designed to explore this possibility.

### Overview of Experiments

Experiments 1 and 2 explore the basic effect of falsely recognizing implied traits under memory instructions. In the second (test) phase, we compared false recognition of implied traits paired with the actual actors’ faces with false recognition of these traits randomly paired with other actors’ faces. If the actual actors’ faces increase false recognition, relative to other familiar actors’ faces, that would support the hypothesis that trait implications (STIs) are linked to the actual actors in explicit memory. Experiment 2 adds novel control traits to the test phase to see whether the mere familiarity of implied traits accounts for some of the false recognition. (Experiments 3 and 4 also include novel control traits.)

The test phase of Experiments 3 and 4 included familiar faces paired with traits opposite in meaning from the implied traits (antonym traits). If people do go beyond the initial trait inference, they should be less likely to falsely recognize antonym traits than control traits. So these studies examine the implications of STIs for other traits.

In Experiment 4 we also sought evidence that pairs of faces and behaviors are processed differently under memory versus impression formation instructions, to further distinguish spontaneous from intentional trait inference processes. We did this by including a condition in which participants formed impressions and by varying the presentation time of sentences and faces in the study phase. We reasoned that additional time would reduce false recognition errors under memory but not impression formation instructions.

Experiments 5 and 6 were designed to see whether the false recognition effects of previous studies depend on explicit memory for the behaviors. We increased the number of behavioral sentences and included measures of participants' ability to retrieve the actors' behaviors (using sentence recognition in Experiment 5 and cued recall in Experiment 6).

Finally, we tested the hypothesis that explicit trait judgments of the actor should predict both false recognition of implied traits and response times (RTs) for correct rejection of implied traits.

### Experiments 1 and 2

Experiments 1 and 2 manipulated the pairing of the faces with the traits implied by the behavioral sentences. Each face was either paired with a trait implied in the sentence presented with the face in the study phase (systematic pairing) or randomly paired with a trait implied in a sentence about another face. If STIs are associated with the person who performed the behavior, participants should be more likely to falsely recognize implied traits when they are paired with that face than when they are randomly paired with a different yet familiar face.

Two things are noteworthy about the experimental paradigm. First, the demands of the recognition task act against the trait inference. Trait inferences lead to incorrect responses. Second, this paradigm rules out explanations of false recognition in terms of mere familiarity. In conditions of both systematic and random pairing of faces and traits, the faces had been presented earlier, and the traits were implied by behavioral sentences. The only difference between these two conditions is the link between the face and the trait-implicating behavior in the study phase. If trait inferences only act as a summary label of behaviors, then the false recognition of such traits should be the same in the systematic and random pairing conditions. However, if the inference is about the person, false recognition should be higher in the systematic condition.

Experiment 2 included an additional condition in which faces from the study phase were paired with control traits unrelated to the behavioral information. Participants should be more likely to falsely recognize implied traits systematically paired with a face than novel, unrelated control traits paired with the same faces. Including novel control traits in this experiment also allowed us to test an auxiliary hypothesis: that merely inferring a trait, whether or not it is linked to an actor, will produce more familiarity and, hence, more false recognition than occurs for novel traits. That is, when participants spontaneously infer traits (even if these traits are

detached from the person who performed the behavior), they may be more likely to falsely recognize the inferred traits than novel traits.

### Method

*Participants.* Twelve undergraduate students from the department of psychology at New York University participated in Experiment 1, and 42 undergraduates participated in Experiment 2 for partial course credit. Participants were randomly assigned to between-subjects replication conditions.

*Stimulus material.* Thirty-six sentences were selected from those collected by Uleman (1988) and his students. These were modified so that the pronouns were replaced by personal names. For 12 of the sentences, the implied trait was included in the behavioral sentence. For example, the sentence, "He threatened to hit her unless she took back what she said," was replaced by the sentence, "Andrew was so aggressive that he threatened to hit her unless she took back what she said." These 12 sentences were used as filler sentences in the experiments. All 36 sentences were paired with photos of faces taken from a yearbook of graduates of Brown University.

To make sure that the sentences in the context of the faces implied the specific trait, we presented the 24 faces and sentences that did not include an implied trait to 24 participants, who rated the person on three 11-point trait scales ranging from 0 (*not at all*) to 10 (*extremely*). For each stimulus pair, participants rated the person on three dimensions: the implied trait, an unrelated trait, and a trait opposite in meaning from the implied trait (an antonym trait). The unrelated traits were used as control traits in Experiments 2, 3, and 4. The antonym traits were used in Experiments 3 and 4. To counterbalance the order of the stimuli, we created 24 different orders, following a Latin square design. Each participant was randomly assigned to 1 of the 24 orders.

The differences in trait ratings were highly significant. Participants rated the target person higher on the implied trait ( $M = 7.58$ ,  $SD = 0.89$ ) than on the unrelated trait ( $M = 5.05$ ,  $SD = 0.88$ ),  $t(23) = 8.58$ ,  $p < .001$ . Participants also rated the target person lower on the antonym trait ( $M = 2.52$ ,  $SD = 1.04$ ) than on the unrelated trait,  $t(23) = 12.80$ ,  $p < .001$ . In fact, for all 24 participants, the ordering of the ratings was the same: highest ratings on the implied trait, intermediate on the unrelated trait, and lowest on the antonym trait. This was also the case for analyses at the level of the stimuli. The differences between implied and unrelated trait ratings were highly significant,  $t(23) = 17.22$ ,  $p < .001$ , as were the differences between unrelated and antonym trait ratings,  $t(23) = 11.16$ ,  $p < .001$ .

*Procedure.* All participants were told that this was a study of how people remember information. Participants worked individually in sound-proof cubicles, and instructions were presented by a computer. Participants were told that the experiment consisted of two parts. In the first part, they would be shown pictures of people with information about these people, and in the second part, their memory would be tested. The type of memory test was not specified in the instruction. The experiment started with a practice trial and, if everything was clear to the participant, continued with the study phase. Each participant was presented with 36 stimulus pairs (trials) of a face and a behavioral sentence. The order of the 36 trials was randomized for each participant by the computer. Each trial (a face plus a sentence) was presented for 10 s. The time delay between trials was 2 s. In 12 of the trials, the sentences contained the trait implied by the behavior. In the remaining 24 trials, the sentences only described behaviors.

After the 36 study trials, participants were told that in the second part, they would be presented with the faces from the first part of the experiment, each accompanied by a single word. Their task was to decide whether they had seen the word in the sentence about the person. The participant's task was to press the *old* key on the keyboard (the *M* key, labeled *old*) if they believed that they had seen the word in the study phase or the *new* key (the *X* key, labeled *new*) if they believed that they had not

seen the word. Participants were asked to work as quickly as possible. To familiarize them with the task, we gave the participants two practice trials. The face and the sentence from the practice trial in the study phase were presented again for 5 s. Then the face was presented with a trait word, which was part of the sentence. Before this trial, participants were reminded about the *old/new* decision. All participants pressed the *old* key. After the trial, participants were told that this was the correct decision and then were presented with a face and a trait word that was not part of the sentence. All participants pressed the *new* key. Participants were told that this was the correct decision, reminded to work as quickly as possible, and asked to continue with the experiment if everything was clear.

In the test phase, participants were presented with the 36 faces from the study phase, each paired with a trait word. The order of the test trials was randomized for each participant by the computer. The trait word was presented below the face. Each test trial stayed on the screen until the participant's response. The next trial followed after 2 s. In 12 trials, participants were presented with the faces from the filler sentences, which contained trait words, paired with the trait words actually presented.

In Experiment 1, the remaining 24 faces were divided into two groups. Within each group of 12 faces, each face was randomly paired with a trait implied in a sentence about another face. Two experimental versions, or replication conditions, were created. In the first version, 12 of the faces were paired with traits implied about them, and 12 were randomly paired with traits implied about a different face. In the second version, the former 12 faces were randomly paired with traits implied about a different face, and the latter 12 were paired with traits implied about them. Participants were randomly assigned to one of the two replication conditions. Thus, the overall design was a mixed 3 (trait: presented vs. implied and systematically paired vs. implied and randomly paired)  $\times$  2 (replication) analysis of variance (ANOVA) with the first factor within subject and the second between subjects.

In Experiment 2, these same 24 faces paired with behavioral sentences were divided into three groups. Within each group, each face was either paired with the implied trait, randomly paired with an implied trait about a different face, or paired with a new, unrelated control trait. Three replication conditions were created. In each one in the test phase, 8 faces were paired with their own implied traits, 8 were paired with implied traits about a different face, and 8 were paired with new, unrelated control traits. The groups of 8 faces were counterbalanced across replication conditions. Thus, each face appeared in the test phase in all three combinations across participants: with its own implied trait, with an implied trait about another face, and with a control trait. The overall design was a 4 (trait: presented vs. implied and systematically paired vs. implied and randomly paired vs. unrelated control)  $\times$  3 (replication) ANOVA, with the first factor within subject and the second between subjects.

**Analyses.** The computer recorded the response to each test trial (*old* vs. *new*) and the RT to the nearest millisecond. The proportion of false recognition of traits and the mean RTs of the correct responses were analyzed at both levels of participants and stimuli.<sup>3</sup> There are two reasons why the RTs analyses were focused on the correct rejection of implied traits. First, the predictions for these RTs are directly derived from the hypotheses, whereas the predictions for RTs on trials with incorrect responses are unclear. Incorrect responses are multiply determined, and there is evidence that the distributions of RTs for error trials and correct trials have different properties (Luce, 1986). Second, for some experimental conditions there were very few incorrect responses, resulting in a large proportion of missing data for analyses across different conditions. In all experiments, RTs below 250 ms were deleted. Such RTs were too short to reflect comprehension of the material. All analyses were conducted on untransformed data.

## Results

**Recognition.** In both experiments, the effects involving replications were not significant. The correct recognition of presented

traits was .93 ( $SD = .09$ ) in Experiment 1 and .83 ( $SD = .12$ ) in Experiment 2. Most important and in support of our hypothesis, in both experiments participants were almost twice as likely to falsely recognize implied traits that were paired with the face of the actor who performed the behavior than to falsely recognize implied traits that were randomly paired with a different face (see Figure 1). The respective proportions were .47 ( $SD = .27$ ) and .26 ( $SD = .18$ ) in Experiment 1,  $t(11) = 2.90$ ,  $p < .015$ ; and .42 ( $SD = .24$ ) and .18 ( $SD = .16$ ) in Experiment 2,  $t(41) = 6.35$ ,  $p < .001$ . This effect was also significant at the level of the stimuli,  $t(23) = 3.37$ ,  $p < .003$ , for Experiment 1, and  $t(23) = 5.26$ ,  $p < .001$ , for Experiment 2. In Experiment 2, participants were also more likely to falsely recognize systematically paired implied traits than novel control traits paired with the same faces ( $M = .15$ ,  $SD = .16$ ),  $t(41) = 7.30$ ,  $p < .0001$  (see Figure 1), as predicted.

The mean proportions of false recognition for randomly paired implied traits and control traits were in the direction predicted by our auxiliary hypothesis: that randomly paired inferred trait rates would be higher, probably because the spontaneous inference makes them familiar. But the difference (.03) did not reach significance either at the level of participants,  $t(41) = 1.62$ ,  $p > .11$ , or at the level of stimuli,  $t(23) = 1.00$ ,  $p > .33$ . However, a closer, post hoc inspection of the control and randomly paired implied traits showed that these traits were mismatched on valence. Whereas all 24 control traits were positive, only 14 of the 24 randomly paired implied traits were positive. Positive traits are more likely to be falsely recognized than are negative traits (e.g., Todorov, 2002; see also the *Effects of Trait Valence on False Recognition* section, below). This difference in the valence of the types of traits might have contributed to the lack of significant differences. In fact, when analyses were restricted to positive traits alone to control for valence, participants were significantly more likely to falsely recognize randomly paired implied traits ( $M = .22$ ,  $SD = .24$ ) than control traits ( $M = .15$ ,  $SD = .16$ ),  $t(41) = 2.33$ ,  $p < .025$ .

**RTs.** In both experiments, the difference between RTs for correct rejection of implied traits systematically and randomly paired with faces was not significant,  $t < 1.00$ .<sup>4</sup> However, in

<sup>3</sup> In all six experiments, the correct recognition of traits presented earlier (the hit rate) was high and significantly different from the false recognition of implied traits. The hit rate is reported for all experiments but is not analyzed in detail, primarily because it is not relevant to the theoretical predictions. The trials with presented traits served as fillers. In addition, they were not manipulated experimentally at the level of stimuli. That is, faces paired with trait sentences in the study phase always appeared with that trait word in the test phase. In contrast, faces paired with trait-implying behavioral sentences appeared either with the implied trait or with some other trait in the test phase. That allowed for the analysis of false recognition at the level of stimuli.

<sup>4</sup> The analysis revealed a main effect of trait,  $F(2, 20) = 3.73$ ,  $p < .042$ , indicating that participants were faster to make a correct decision about presented traits than about implied traits. Such effects were also found in the other experiments. We do not report these effects because they are not theoretically relevant and their interpretation is further complicated by the fact that the correct responses for presented and implied traits were confounded with the response keys. Correct responses for presented traits were made with the *M* key (right hand), whereas correct responses for implied traits were made with the *X* key (left hand).

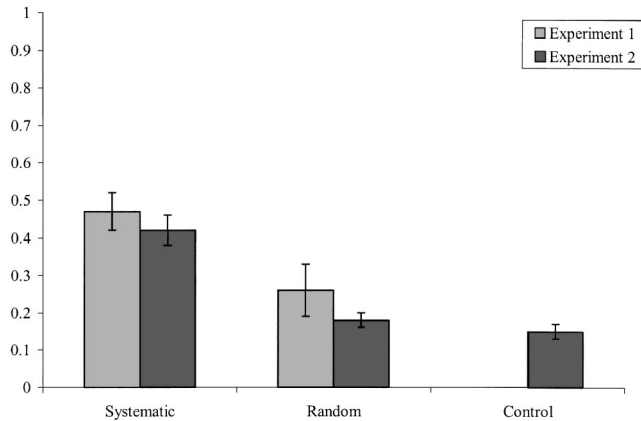


Figure 1. Mean proportion of false recognition of traits as a function of face–trait pairing (Experiments 1 and 2). In the systematic condition, the implied trait was paired with the actor's face. In the random condition, the implied trait was randomly paired with a familiar face of another actor. In the control condition, the actor's face was paired with an unrelated trait.

Experiment 2, correct rejection of systematically paired implied traits was slower ( $M = 2,765.00$ ,  $SD = 1,105.00$ ) than correct rejection of control traits paired with the same face ( $M = 2,457.00$ ,  $SD = 769.00$ ),  $t(41) = 3.53$ ,  $p < .001$ . Correct rejection of randomly paired implied traits ( $M = 2,736.00$ ,  $SD = 1,044.00$ ) was also significantly slower than the rejection of control traits,  $t(41) = 2.88$ ,  $p < .006$ . The same effects were revealed at the level of stimuli,  $t(23) = 3.46$ ,  $p < .002$ , for the difference between RTs to systematically paired implied traits and control traits, and  $t(23) = 3.47$ ,  $p < .002$ , for the difference between randomly paired implied traits and control traits.

## Discussion

Both of our hypotheses were supported, and there were other interesting effects. The recognition findings show that participants spontaneously inferred traits and that these traits were associated with the person who performed the behavior. RT findings suggest that participants inferred the implied trait when presented with the behavior and that the resulting trait familiarity slowed the correct response of rejecting the implied trait, regardless of whether it was systematically or randomly paired with an actor. Similar familiarity effects were also revealed by post hoc analyses of the false recognition data, with the valence of the traits controlled for. Even when the implied traits were not paired with the relevant actors, participants were more likely to falsely recognize these traits than novel traits.

The lack of significant differences between RTs to systematically and randomly paired implied traits, coupled with the significant differences between false recognition rates for these traits, reveals an interesting dissociation. RTs were more sensitive to trait familiarity produced by prior inferences than were recognition rates. Recognition rates were more sensitive to pairings with relevant actors than to mere familiarity produced by prior inferences. This suggests a context-checking process in which any familiarity from prior inferences slowed participants on recognition trials. Then the presence of the relevant actor provided a

misleading context, because trait and actor were bound together in LTM, raising false recognition rates. When familiarity was low, this context checking was omitted.

## Experiment 3

Experiments 1 and 2 show that spontaneously inferred traits are associated with the person who performed the behavior. Experiment 3 addresses an additional question: Can people use this initial trait inference to go beyond it, even though the inference is implicit? When people have explicit trait knowledge about a person, they clearly can use this knowledge as a basis for other trait inferences (e.g., Rosenberg & Sedlak, 1972; Schneider, 1973). If spontaneously inferred traits are encoded in the person representation, they may also have implications for other traits. Experiment 3 included test trials in which the actors' faces were paired with traits opposite in meaning from the implied traits (antonym traits). If the implied trait is encoded into the person representation and has implications for other traits, participants' false recognition of antonym traits should be lower than the false recognition of control traits unrelated to the person's behavior. The control traits were the novel traits used in Experiment 2. Such traits provide a more appropriate control comparison than do randomly paired implied traits because both these traits and the antonym traits are novel with respect to the behaviors.

## Method

**Participants.** Twenty-four undergraduates from the department of psychology at New York University participated in the study for partial course credit. They were randomly assigned to three between-subjects replication conditions.

**Procedure.** The procedures were the same as in Experiments 1 and 2. The only difference was that the test trials included traits opposite in meaning from the implied traits (antonym traits). The 24 faces that were paired with behavioral sentences were divided into three groups. Within each group of 8 faces, each face was paired with either the implied trait; a new, unrelated control trait; or an antonym trait. As in Experiment 2, the groups of 8 faces were counterbalanced across three replication conditions. The overall design was a 4 (trait: presented vs. implied vs. control vs. antonym)  $\times$  3 (replication) mixed ANOVA with the first factor within subject and the second between subjects.

## Results

**Recognition.** The proportion of correct recognition of presented traits was .81 ( $SD = .11$ ). The effects involving replication were not significant. Participants were more likely to falsely recognize implied traits ( $M = .35$ ,  $SD = .26$ ) than control traits ( $M = .19$ ,  $SD = .21$ ; see Figure 2),  $t(23) = 2.63$ ,  $p < .015$ .

More important, as predicted, participants were less likely to falsely recognize antonyms of the implied traits ( $M = .10$ ,  $SD = .12$ ) than control traits,  $t(23) = 2.67$ ,  $p < .014$ . These effects were also significant at the level of the stimuli,  $t(23) = 3.77$ ,  $p < .001$ , for implied versus control traits, and  $t(23) = 2.25$ ,  $p < .035$ , for antonym versus control traits.

**RTs.** Participants were slower to correctly reject implied traits ( $M = 2,287.00$ ,  $SD = 885.00$ ) than control traits ( $M = 2,175.00$ ,  $SD = 793.00$ ) or antonym traits ( $M = 2,149.00$ ,  $SD = 799.00$ ), but these effects were not significant,  $ts(23) < 1.23$ ,  $ps > .23$ .

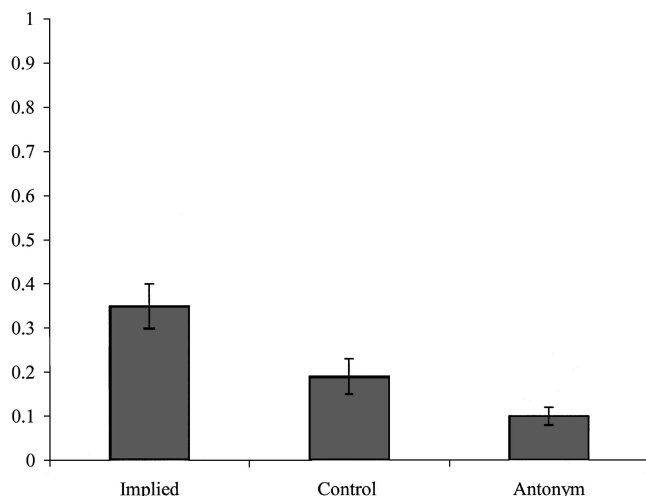


Figure 2. Mean proportion of false recognition of traits as a function of type of trait (Experiment 3).

### Discussion

As in Experiment 2, participants were more likely to falsely recognize implied traits paired with the face of the person who performed the behavior than control traits paired with the same face. More important, Experiment 3 shows that people can go beyond that initial, spontaneous inference about the actor. Participants were less likely to falsely recognize antonym traits of the implied traits than control traits. This finding shows that this new person knowledge had implications for other traits.

The difference between RTs for correct rejection of implied traits and control traits was consistent with the difference found in Experiment 2. Participants were slower in rejecting implied traits than control traits. However, this difference did not reach significance. One reason may be the small sample size ( $N = 24$  rather than 42) and the resulting lack of power in this experiment. Experiment 4 uses a larger sample and includes a replication of Experiment 3.

### Experiment 4

In all previous experiments, participants were given memory instructions. However, it is possible that the presentation of faces paired with behaviors might have triggered impression formation goals. Thus, it is important to demonstrate that memory and impression instructions engage different processes. There is ample evidence in the person perception literature that lists of traits are processed differently under memory and impression goals. For example, in their classic study, Hamilton, Katz, and Leirer (1980) found that impression formation instructions produced better recall and more clustering of lists of trait terms about the same target. And there is clear evidence that a variety of processing goals affect the likelihood of STIs (Uleman & Moskowitz, 1994).

However, Carlston and Skowronski (1994) did not find differences between impression instructions and other instructions, such as familiarization with the material and memorization, in five experiments using the savings paradigm. This was the case even when the impression instructions were trait focused (e.g., "Think

about a specific trait that would describe the person's personality"). The lack of differences in the savings paradigm suggests either that faces paired with paragraphs triggered impression formation goals or that the effect of associating inferred traits with the actor is fairly robust and relatively independent of processing objectives. But even if the effect is robust, it should still be possible to demonstrate dissociations between memory and impression instructions as a function of relevant factors. The purpose of Experiment 4 was to find such a dissociation.

One relevant factor that can produce differential effects on memory and impression instructions is the exposure time to the faces paired with behaviors. Previous on-line evidence for STIs (Uleman et al., 1996; Van Overwalle, Drenth, & Marsman, 1999) suggests that trait inferences from single behaviors occur very rapidly, easily in less than 5 s. If participants were given more encoding time, they could use the extra time to form a more detailed representation of the behavior under memory instructions. Thus, they should be more likely to discriminate between implied and presented information. Under impression instructions, participants do not have a goal of discriminating implied and presented information. In fact, these instructions require participants to engage in explicit inferences, so the extra encoding time should not result in better discrimination between implied and explicit information.

To test these hypotheses, we gave participants either memory or impression instructions, and each face-behavior pair was presented for either 5 s or 10 s. Increasing the time exposure should result in a lower rate of false recognition of implied traits under memory instructions but should not affect this rate under impression instructions.

### Method

**Participants.** One hundred eighty-six undergraduate students from the department of psychology at New York University participated in the study for partial course credit. Participants were randomly assigned to 12 between-subjects conditions.

**Procedure.** The procedures were the same as in Experiment 3. However, in addition to the manipulations in Experiment 3, Experiment 4 manipulated instructions to participants and the presentation time of the faces and behavioral information. Participants were either given a memory instruction, as in Experiments 1, 2, and 3, or an impression formation instruction. In the impression formation condition, participants were told that the study was about how people form impressions of other people and how they determine the personalities of those people. Participants were asked to read the behavioral information and to think about the character of the person who was described in the behavioral sentence. They were told that in the second part of the study they would be asked questions about the character of this person. In addition, half the participants saw the stimuli for 5 s, whereas the other half saw them for 10 s each. The overall design was a 4 (trait: presented vs. implied vs. new control vs. antonym)  $\times$  3 (replication)  $\times$  2 (instruction: impression vs. memory)  $\times$  2 (presentation time: 10 s vs. 5 s) mixed ANOVA with the first factor within subject and the latter three factors between subjects.

### Results

**Recognition.** Effects involving replication were not significant. The ANOVA revealed a significant Trait  $\times$  Presentation Time interaction,  $F(3, 522) = 2.62, p < .05$ , and a significant Trait  $\times$  Instruction  $\times$  Time interaction,  $F(3, 522) = 3.06, p <$

.028. To understand the triple interaction, we performed separate analyses for the impression and memory participants. Under memory instruction, the Trait  $\times$  Time interaction was significant,  $F(3, 273) = 5.98, p < .001$  (see Figure 3). Simple comparisons showed that presentation time only affected the false recognition of implied traits, which were more likely to be falsely recognized under 5-s than under 10-s presentation,  $t(91) = 2.72, p < .008$  (see Table 1). Under impression instruction, the interaction of trait and presentation time was not significant,  $F < 1.00$ . That is, recognition under impression instructions was unaffected by presentation time.

To test the hypotheses that participants are more likely to falsely recognize implied than control traits and less likely to falsely recognize antonym than control traits, we conducted two analyses. A 2 (trait: implied vs. control)  $\times$  2 (instruction)  $\times$  2 (presentation time) mixed ANOVA revealed that participants were more likely to falsely recognize implied traits ( $M = .41, SD = .28$ ) than control traits ( $M = .16, SD = .16$ ),  $F(1, 182) = 171.51, p < .001$ . A triple interaction of trait, instruction and time was marginally significant,  $F(1, 182) = 3.57, p < .06$ . This interaction was described in the previous paragraph, and it did not substantially modify the main effect of trait. Nevertheless, we compared false recognition of implied and control traits at all four combinations of instruction and presentation time. These four simple effects were all significant at the level of participants,  $ps < .001$ , as well as at the level of stimuli,  $ps < .001$ .

A 2 (trait: antonym vs. control)  $\times$  2 (instruction)  $\times$  2 (presentation time) mixed ANOVA only revealed a significant main effect of trait,  $F(1, 182) = 11.72, p < .001$ . Participants were less likely to falsely recognize antonym traits ( $M = .12, SD = .15$ ) than control traits ( $M = .16, SD = .16$ ). This was also the case at the level of stimuli,  $F(1, 23) = 4.27, p < .05$ .

**RTs.** The analysis of RTs for correct responses at the level of participants revealed a significant main effect of trait,  $F(2, 340) = 16.09, p < .001$ , and a significant Trait  $\times$  Replication interaction,  $F(4, 340) = 11.21, p < .001$ . In one of the replication conditions, the RTs for rejecting implied traits did not differ from RTs for rejecting control traits, whereas in the other two replications, rejecting implied traits took longer than rejecting control

Table 1

*False Recognition of Traits as a Function of Type of Trait, Instruction, and Presentation Time of Faces and Behavioral Information (Experiment 4)*

| Condition and statistic | Hits | False recognition of traits |           |                     |
|-------------------------|------|-----------------------------|-----------|---------------------|
|                         |      | Implied                     | Unrelated | Opposite to implied |
| Memory 10 s             |      |                             |           |                     |
| <i>M</i>                | .86  | .32                         | .14       | .11                 |
| <i>SD</i>               | .11  | .28                         | .15       | .15                 |
| Memory 5 s              |      |                             |           |                     |
| <i>M</i>                | .82  | .47                         | .18       | .14                 |
| <i>SD</i>               | .11  | .26                         | .17       | .14                 |
| Impression 10 s         |      |                             |           |                     |
| <i>M</i>                | .82  | .41                         | .15       | .14                 |
| <i>SD</i>               | .14  | .29                         | .15       | .15                 |
| Impression 5 s          |      |                             |           |                     |
| <i>M</i>                | .84  | .42                         | .18       | .11                 |
| <i>SD</i>               | .12  | .27                         | .18       | .16                 |

traits. Because this interaction is theoretically irrelevant and involved only one of the replications, we compared the response times across replications. The correct rejection of implied traits ( $M = 2,544.00, SD = 1,290.00$ ) was slower than the rejection of control traits ( $M = 2,251.00, SD = 725.00$ ),  $t(181) = 3.85, p < .001$ , and the rejection of antonym traits ( $M = 2,278.00, SD = 811.00$ ),  $t(181) = 3.68, p < .001$ . Both effects were also significant at the level of stimuli,  $t(23) = 2.69, p < .013$ , for implied versus control traits, and  $t(23) = 2.19, p < .039$ , for implied versus antonym traits.

The analysis at the level of stimuli also found main effects for instruction,  $F(1, 23) = 5.53, p < .028$ , and for presentation time,  $F(1, 23) = 6.23, p < .02$ . Responses were faster under impression instruction ( $M = 2,279.00, SD = 193.00$ ) than memory instruction ( $M = 2,397.00, SD = 213.00$ ). Responses were also faster after the 10-s initial presentation ( $M = 2,268.00, SD = 206.00$ ) than after the 5-s initial presentation ( $M = 2,408.00, SD = 217.00$ ).

## Discussion

Experiment 4 replicates and extends the findings of Experiment 3. Most important, this experiment demonstrates that encoding time had different effects on memory and impression participants. In the memory condition, false recognition of implied traits was higher when the stimuli were presented for 5 s than when the stimuli were presented for 10 s. It is important to note that the presentation time manipulation only affected the false recognition of implied traits, not unrelated or antonym traits. In contrast, the encoding time did not affect the false recognition of implied traits under impression instructions.

Presentation time also affected RTs. They were slower after 5-s than 10-s presentations, regardless of instructions, suggesting greater decision uncertainty with shorter presentations. It is interesting that this greater uncertainty was only warranted under memory instructions for recognition of implied traits, where false recognition was in fact higher. Neither hits nor false recognition differed under other conditions.

This suggests that different processes were involved under memory than impression formation instructions, even though in-

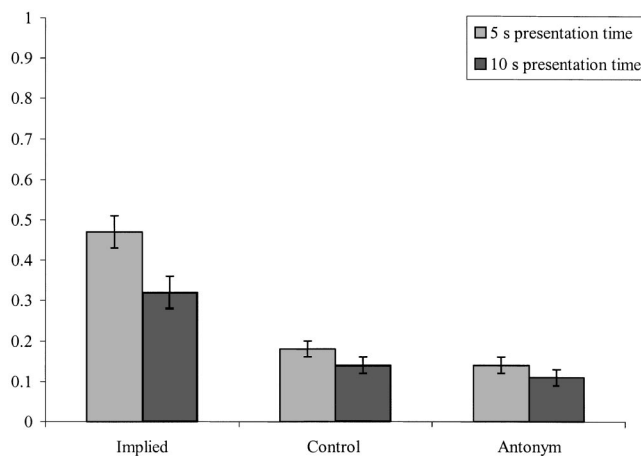


Figure 3. Mean proportion of false recognition of traits as a function of type of trait and presentation time of faces with behavioral sentences under memory instruction (Experiment 4).

ferences of actors' traits occurred under both instructions within 5 s. Participants anticipating a test of their memory seem to have used their extra 5 s to attend more carefully to what was actually stated rather than merely implied, whereas those forming impressions did not. After all, their task was to make inferences and form impressions, not commit details to memory, and an additional 5 s made no difference in their inferences. Then, when faced with a recognition memory task, all participants were more cautious (slower) when their exposure had been brief. But taking longer did not increase accuracy across the board. The only accuracy difference was for implied traits under memory instructions. That is, the longer RTs after 5-s presentations had no general effect on accuracy. What was critical for accuracy was how participants used their time at encoding.

As did Experiment 3, this experiment failed to find facilitation of RTs for antonyms, in spite of the fact that false recognition of antonyms was reliably lower in both studies. This finding suggests that only implied traits, not antonyms, were inferred at encoding. Antonym probes may be compared with the person representation at test (like all other probes) and, because they are more clearly or sharply contradicted by that representation than control traits are, they may be rejected with more certainty (less false recognition) but no more speed than are control words. This would be consistent with Skowronski and Shook's (1997) conclusion that impressions do not invariably include information about antonyms. But the lower false recognition rate does show that antonym-relevant information is clearer in these spontaneous representations than is information relevant to other traits.

### Experiment 5

We have argued that participants inferred the trait during the behavioral presentation and that this trait was encoded as part of the actor representation. Then, during the recognition test, participants failed to discriminate between the inference and the actual information. There is an alternative explanation of the high false recognition of implied traits in terms of retrieval processes. According to this explanation, participants did not infer or encode any traits about the actor. When presented with the trait word in the recognition task, they recalled the actor's behavior and then checked the plausibility of the trait word against the recalled behavior. Implied traits fit the behavioral sentences better than did randomly paired implied traits (Experiment 1 and 2) or unrelated traits (Experiment 3 and 4), producing the higher false recognition of implied traits. The purpose of Experiment 5 (and 6) was to look for our basic false recognition effect under conditions that rule out this retrieval explanation.

Participants were presented with 120 faces paired with behavioral sentences for 5 s per stimulus pair. Replicating the false recognition findings with such a large number of stimuli would be, by itself, an argument against the retrieval explanation, because this large number of sentences should make retrieval of any particular sentence more difficult. We also included a sentence recognition test at the end of the experiment. Each face that was paired with a trait word earlier in the experiment was presented with four sentences—the actual sentence, and three foils. According to the retrieval explanation, the false recognition findings should be obtained only for sentences recognized by participants. Alternatively, if the false recognition findings are driven by en-

coding processes, the effect should be obtained for both recognized and unrecognized sentences.

Possible familiarity and trait inference effects on sentence recognition were controlled in that three out of the four sentences in the latter test implied the same trait, and two of them were presented in the study phase of the experiment. These foils were necessary because if participants inferred a trait during encoding and only one of the sentences implied this trait, they could use the trait inference to guess the correct sentence. Similarly, if only one of the sentences had been presented in the study phase, participants could rely on simple familiarity to guess the correct sentence.

Because each trait was implied by two different sentences paired with different faces in the study phase, only half of the faces were paired with trait words in the word recognition test phase, so that the same trait would not be presented twice in the test phase. The other half of the faces were paired with nontrait words. These nontrait words included personal names, words presented in the sentences, and words implied but not presented in the sentences. For example, for the sentence "Glen phoned for help while the others just screamed," the name was *Glen*, the presented nontrait word was *help*, and the implied nontrait word was *danger*.

This word pairing has the advantage over previous studies of not focusing the participants' attention exclusively on trait words in the word recognition test phase. It also allowed us to explore an additional hypothesis. We were interested in whether the individuation of trait inferences is comparable to the individuation of personal attributes that are actually presented, such as names. For this purpose, some of the names were paired with the faces with which they were presented in the study phase, and some were randomly paired with different faces.

### Method

**Participants.** Twenty-four undergraduate students from the department of psychology at New York University participated in the study for partial course credit. Participants were randomly assigned to two between-subjects conditions.

**Procedure.** The instructions were the same as in Experiments 1, 2, and 3. In the study phase, participants were presented with 120 faces paired with behavioral sentences. Each face-sentence pair was presented for 5 s, and the intertrial interval was 2 s. The order of trials was randomized for each participant. Forty of the sentences contained explicit traits. Thus, the ratio of sentences with implied traits and sentences with explicit traits (2:1) was the same as in all previous experiments. In the word recognition task, participants made recognition decisions for 120 words paired with faces. The intertrial interval was 2 s, and the order of trials was randomized for each participant.

The 120 different behaviors consisted of 60 pairs of behaviors, each pair implying the same trait. For example, the trait *honest* was implied by the following sentences: "Bonnie told the cashier that she got too much change," and "Emily returned the lost wallet with all the money in it." In the word recognition test phase, if one of the faces (*Bonnie* or *Emily*) was paired with the trait word (*honest*), the other face was paired with a nontrait word. Thus, 60 of the faces were paired with trait words, and 60 were paired with nontrait words. In the group of 60 face-trait pairs, 20 of the faces were paired with actually presented trait words. The other 40 faces were randomly divided into two groups. In keeping with the procedures of Experiment 1, within each group each face was randomly paired with an implied trait about a different face. Two replication conditions were created in which the systematically and randomly paired faces with traits were counterbalanced. In the group of 60 face-nontrait-word pairs, 20 of the faces were paired with nontrait words that were part of the sentence, 20

were paired with implied nontrait words, 10 were paired with the actor's name, and 10 were randomly paired with a name of a different actor. Faces paired with implied and presented nontrait words were counterbalanced in the replication conditions, as were faces systematically and randomly paired with personal names.

After the word recognition test, participants were presented with a sentence recognition test for all faces that were paired with implied traits (both systematically and randomly) in the word recognition test phase ( $n = 40$ ). The order of trials was randomized for each participant, and the intertrial interval was 2 s. Each face was presented with four sentences. One of the sentences was the sentence presented with the face (e.g., "Bonnie told the cashier that she got too much change"). Another sentence implied the same trait but had been paired with another face in the study phase (e.g., "Emily returned the lost wallet with all the money in it"). The only difference was that for the sentence recognition test, the name of the relevant actor was changed to the name of the test actor (*Emily* to *Bonnie*). The remaining two sentences were new, one implying the same trait, and another containing most of the words of the actual sentence but not implying the trait. All sentences started with the name of the actor.

## Results

The correct recognition of presented traits was .68 ( $SD = .14$ ).<sup>5</sup> Across both replications, participants were more likely to falsely recognize systematically paired implied traits ( $M = .42$ ,  $SD = .17$ ) than randomly paired implied traits ( $M = .34$ ,  $SD = .17$ ),  $t(23) = 2.85$ ,  $p < .009$  (see Figure 4). This effect was also significant at the level of the stimuli,  $t(39) = 2.23$ ,  $p < .032$ . Most important, the effect did not depend on the sentence recognition. The respective proportions of false recognition of systematically and randomly paired implied traits were .42 and .33 for unrecognized sentences and .42 and .35 for recognized sentences. On average, participants recognized 62% of the sentences.

The correct recognition of personal names ( $M = .55$ ,  $SD = .19$ ) was higher than the false recognition of randomly paired names ( $M = .47$ ,  $SD = .20$ ), but the difference only approached significance,  $t(23) = 1.75$ ,  $p < .095$ . Although the interaction of pairing (systematic vs. random) and word (implied trait vs. name) was not significant, the effect size for false recognition of implied traits ( $r = .51$ ) was larger than the effect size for recognition of names

( $r = .34$ ). This suggests that systematic versus random pairing had a larger effect on recognition of implied traits than on presented names.

As in Experiments 1 and 2, the response times for correct rejection of systematically and randomly paired implied traits did not differ significantly,  $t < 1.00$ .

## Discussion

Experiment 5 replicates the findings of Experiments 1 and 2 with a large number of behaviors, each presented only for 5 s. Most important, this experiment establishes that the effect was not dependent on the retrieval of the actor's behavior. Whether or not participants were able to recognize the behavioral sentence, they were more likely to falsely recognize traits systematically paired with the actor's face than randomly paired implied traits. This finding provides support for the hypothesis that traits are inferred and encoded as part of the actor representation during the behavioral presentation. It also suggests that sentence recall played little or no significant role in the false recognition decisions.

This experiment also shows that the individuation of implied traits is similar, if not larger, in size to the individuation of explicitly presented personal attributes, such as names.

## Experiment 6

A main assumption of the retrieval explanation of our false recognition findings is that participants retrieve the actor's behavior for the recognition of trait words in the test phase. The findings from the sentence recognition task in Experiment 5 were inconsistent with this explanation. In Experiment 6 we used a cued recall test to test the retrieval explanation more directly. After the word recognition test phase, participants were presented with the faces that had been paired with implied trait words and asked to recall the sentences about these faces. Given the large number of behaviors and their short presentation times, we expected the recall of behaviors to be poor. More important, we expected the effect of false recognition of implied traits to be independent of the recall of the actors' behaviors.

We also used this experiment to explore another issue: How easily do names become associated with actors' faces, relative to other explicit and inferred attributes? Our hypothesis was that names would show smaller false recognition effects than other attributes because, unlike other attributes, they were associated with faces in an entirely arbitrary and uninformative manner.

## Method

**Participants.** Twenty undergraduate students from the department of psychology at New York University volunteered and were paid for their participation. Participants were randomly assigned to two between-subjects conditions.

**Procedure.** The procedures were the same as in Experiment 5, except that (a) each type of stimulus word (20 presented traits, 40 implied

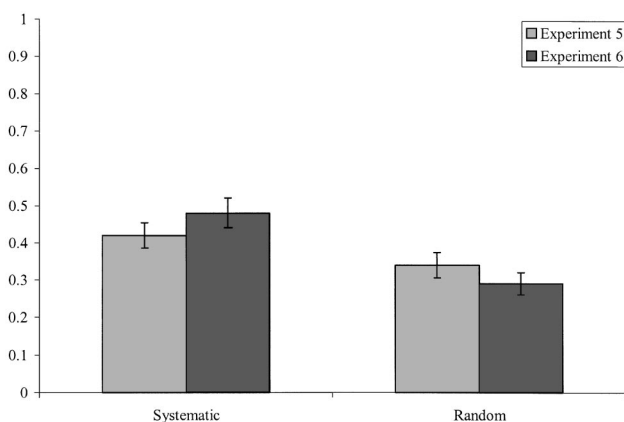


Figure 4. Mean proportion of false recognition of traits as a function of face-trait pairing (Experiments 5 and 6). In the systematic condition, the implied trait was paired with the actor's face. In the random condition, the implied trait was randomly paired with a familiar face of another actor.

<sup>5</sup> The correct recognition of nontrait words was .57 ( $SD = .10$ ). This proportion was higher than the false recognition of implied nontrait words ( $M = .45$ ,  $SD = .16$ ),  $t(23) = 3.71$ ,  $p < .001$ . Because the trait and nontrait words differ as word types (adjectives vs. nouns in most cases) and in frequency, we did not compare them directly.

traits, 20 presented nontrait words, 20 implied nontrait words, and 20 names) was systematically and randomly paired and (b) participants were presented with a cued recall test rather than with a sentence recognition test. Two replication conditions were created in which the systematically and randomly paired words were counterbalanced.

After the word recognition test phase, participants were presented with a cued recall test for all faces that had been paired with implied traits (both systematically and randomly) in the word recognition test ( $n = 40$ ). The order of trials was randomized for each participant, and the intertrial interval was 2 s. Each trial presented a face, and participants were instructed to write down the sentence that had been presented with the face.

Two judges who were unaware of the replication conditions coded participants' recall into two exhaustive categories: correct recall of actor's behavior versus lack of recall or incorrect recall. Behaviors were coded as correct in terms of their meaning. For example, recall of the behavior "she did not go out to study" was coded as correct recall of the sentence "Elena turned down three parties to study for organic chemistry." The agreement between the judges was 95%. Cases that were coded differently were resolved by a third judge.

## Results

Across both replications, participants were more likely to falsely recognize implied traits systematically paired with the actors' faces ( $M = .48$ ,  $SD = .18$ ) than implied traits randomly paired with familiar faces ( $M = .29$ ,  $SD = .15$ ),  $t(19) = 4.56$ ,  $p < .001$  (see Figure 4). This effect was also significant at the level of stimuli,  $t(39) = 4.24$ ,  $p < .001$ , replicating our basic finding from previous studies.

How was behavior recall related to false recognition of traits? Participants correctly recalled 12% of the behaviors. In those cases, they were more likely to falsely recognize systematically paired implied traits ( $M = .49$ ,  $SD = .39$ ) than randomly paired traits ( $M = .18$ ,  $SD = .34$ ),  $t(19) = 2.45$ ,  $p < .024$ . More important, the respective proportions for the majority of cases, in which participants did not recall the behavior, were as follows: systematically paired implied traits,  $M = .45$ ,  $SD = .19$ ; and randomly paired traits,  $M = .29$ ,  $SD = .14$ ;  $t(19) = 3.70$ ,  $p < .002$ . Although the difference between the proportions of systematically and randomly paired traits seems to suggest that the false recognition effect is stronger for recalled behaviors, the interaction effect of type of trait and recall of behavior was not significant,  $F(1, 19) = 1.77$ ,  $p > .20$ . Moreover, analyses of the effect sizes indicated that the effect size for nonrecalled behavior ( $r = .65$ ) was actually larger than the effect size for recalled behaviors ( $r = .49$ ) because of the large standard errors in the latter condition. These results are inconsistent with a retrieval account of our false recognition findings.

How were other attributes related to actors in memory? The correct recognition of systematically paired explicit traits ( $M = .68$ ,  $SD = .16$ ) was significantly higher than the false recognition of randomly paired explicit traits ( $M = .32$ ,  $SD = .17$ ),  $t(19) = 6.99$ ,  $p < .001$ . Similarly, the correct recognition of systematically paired explicit nontrait words ( $M = .56$ ,  $SD = .14$ ) was higher than the false recognition of randomly paired nontrait words ( $M = .30$ ,  $SD = .17$ ),  $t(19) = 6.11$ ,  $p < .001$ . And the false recognition of systematically paired implied nontrait words ( $M = .40$ ,  $SD = .15$ ) was higher than the false recognition of randomly paired implied nontrait words ( $M = .27$ ,  $SD = .18$ ),  $t(19) = 2.68$ ,  $p < .015$ . In contrast to these findings, the correct recognition of personal names ( $M = .50$ ,  $SD = .14$ ) was practically the same as

the false recognition of names randomly paired with faces ( $M = .49$ ,  $SD = .25$ ),  $t < 1.00$ . So the evidence is, again, that names were less individuated than other personal attributes.

The RTs for correct rejection of systematically and randomly paired implied traits were not significantly different,  $t < 1.00$ .

## Discussion

This experiment provides additional evidence that the effect of high false recognition of implied traits is not dependent on the retrieval of the actor's behavior. These findings strongly suggest that trait inferences were made during the encoding of the behavior.<sup>6</sup> The finding that the false recognition of systematically paired implied nontrait words was higher than the false recognition of randomly paired nontrait words also suggests that participants engaged in a variety of other inferences during the presentation of the actor's behavior.

Consistent with Experiment 5, Experiment 6 also shows that inferred traits are more strongly individuated than are personal names. Although names are an important part of person identity, previous studies have shown that they are less likely to be learned and that they are retrieved more slowly than other person attributes, such as occupations and nationality (Johnston & Bruce, 1990; McWeeny, Young, Hay, & Ellis, 1987). Experiments 5 and 6 extend these findings by showing that inferred person attributes can be more strongly associated with persons than can explicitly presented names.

## Analyses of Effect Sizes Across Experiments

In this section, we report the effect sizes across experiments at the levels of participants and stimuli for false recognition of implied traits, false recognition of antonym traits, and RTs for correct rejection of implied traits. We focus only on memory conditions here because in these conditions, the trait inferences are spontaneous and all experiments included such conditions. Effect sizes were calculated in terms of Pearson  $r$ . We calculated average effect sizes by transforming the experimental effect sizes into  $z$ , and weighting them by their respective sample sizes.

For the false recognition of person-associated implied traits, the average effect size across the six experiments was .66 at the level of participants and .63 at the level of stimuli (Table 2). (By

<sup>6</sup> One might suggest (as one reviewer did) that a fairer test of the retrieval explanation would be an examination of sentence recall when both the face and the trait are present as cues, because this would resemble the test phase more precisely. That is, even though false recognition effects occurred in Experiment 6 without sentence recall with photo cues, adding trait cues might reveal much higher sentence recall and thereby account for the false recognition effects. There are two reasons we did not run such a study. First, with each photo in Experiment 5 there were three sentences implying the same trait. Although not equivalent to presenting the trait itself along with the photo, these strong trait implications did not reveal sentence recognition levels that would account for our false trait recognition effects. Second and more important, if adding traits to the recall cues of Experiment 6 were to increase sentence recall, we would interpret this as evidence for encoding specificity (Tulving & Thomson, 1973)—that is, as evidence that the trait had been inferred at encoding and thereby became an effective retrieval cue. Thus, such a study would be (at best) inconclusive in choosing between encoding and retrieval explanations.

Table 2  
*Effect Sizes ( $r$ ) of False Recognition of Person-Associated Traits as a Function of Experiment and Type of Trait at the Levels of Participants and Stimuli*

| Experiment          | Effect sizes |         |         |         |
|---------------------|--------------|---------|---------|---------|
|                     | Participants |         | Stimuli |         |
|                     | Implied      | Antonym | Implied | Antonym |
| Experiment 1        | .66          |         | .58     |         |
| Experiment 2        | .70          |         | .74     |         |
| Experiment 3        | .48          | .49     | .62     | .42     |
| Experiment 4 (10 s) | .57          | .23     | .76     | .27     |
| Experiment 4 (5 s)  | .79          | .26     | .85     | .31     |
| Experiment 5        | .51          |         | .34     |         |
| Experiment 6        | .72          |         | .56     |         |
| Average effect size | .66          | .30     | .63     | .33     |

comparison, the average effect size for the impression condition from Experiment 4 was .69 at the level of participants and .88 at the level of stimuli.) In Experiments 1, 2, 5, and 6, the effect size was estimated from the difference between false recognition of implied traits paired with the actors and false recognition of implied traits randomly paired with other persons. In Experiments 3 and 4, effect sizes were estimated from the difference between false recognition of implied traits paired with the actors and false recognition of control traits paired with the same persons.

Assuming that the first trait inference is based on the behavior and that this inference has implications for other traits, such as antonyms, one would expect a weaker effect for false recognition of antonym traits than for implied traits. As shown in Table 2, effect sizes for the lower rate of false recognition of antonyms to implied traits relative to control traits were smaller than the effects for false recognition of implied traits. The average effect sizes across Experiments 3 and 4 were .30 at the level of participants and .33 at the level of stimuli.

In Experiments 2, 3, and 4, RTs for correct rejection of implied traits were longer than RTs for correct rejection of control traits. The RT findings were generally less consistent than the false recognition findings. However, as shown in Table 3, the differences in response times across experiments were consistent. These effect sizes were .29 at the level of participants and .36 at the level of stimuli.

#### Actor-Linked STIs: Correlation Analyses at Level of Stimuli

If the recognition responses reflect STIs about the person, explicit trait judgments of the person should predict both the false recognition of implied traits and the RTs for correct rejection of implied traits. The stronger the inference about the person is, the more familiar the trait word in the context of the person should be and, hence, the harder and slower the correct rejection of the trait should be.

To test these hypotheses, we conducted correlational analyses at the level of stimuli.<sup>7</sup> As in the analysis of the effect sizes and for the same reasons, this analysis was on data from the memory conditions. Experiments 1, 2, 3, and 4 used the same 24 pairs of

faces and behavioral sentences in the study phase of the experiments. In the test phase, each face was paired with either an implied trait, a control trait, or an antonym of the implied trait, which thus created 72 face–trait pairs. Aggregating across the experiments made it possible to obtain reliable estimates of the mean proportion of false recognition and the mean RTs for each face–trait pair in the memory conditions. Estimates for the implied traits were based on data from 57 to 60 participants per face–trait pair, estimates for control traits were based on data from 51 to 54 participants per face–trait pair, and estimates for antonym traits were based on data from 37 to 40 participants per face–trait pair. These participant numbers differ because of different experimental conditions in each experiment.

As outlined in the *Stimulus material* section of Experiment 1, the 24 stimulus pairs of faces and behaviors were rated for their trait implications. Specifically, participants who did not participate in any of the experiments rated each of the 24 actors on the implied trait, the unrelated trait, and the antonym trait, providing 72 ratings. Note that these ratings were explicit, intentional judgments of the actors. Across types of traits, the mean ratings correlated highly with both the rate of false recognition,  $r(72) = .81$ ,  $p < .001$ , and the correct RTs,  $r(72) = .44$ ,  $p < .001$ . The stronger the trait judgment about the actor was, the higher the false recognition was, and the slower the RTs for correctly rejecting the traits were. The false recognition rate and the RTs were also correlated,  $r(72) = .56$ ,  $p < .001$ .

Undoubtedly, the high correlation between person ratings and false recognition was due to our inclusion of a wide range of traits—from traits specifically implied by the behavioral sentence to traits opposite in meaning. Because the types of traits differed, we conducted separate analyses for each type of trait. These analyses revealed different patterns for implied traits and the other types of traits. As shown in Table 4, the explicit person judgments significantly predicted both the false recognition and the RTs for correct rejection of implied traits. However, the RTs and the false recognition did not correlate significantly. The pattern remained the same after the analyses controlled for the effect of the third variable on the correlated variables. In contrast to the pattern for implied traits, the explicit judgments predicted neither the false recognition nor the RTs for both unrelated and antonym traits. But the false recognition and the RTs correlated significantly:  $r(24) = .45$ ,  $p < .026$  (partial  $r = .46$ , controlling for the effect of judgments) for unrelated traits, and  $r(24) = .50$ ,  $p < .013$  (partial  $r = .51$ ) for antonym traits.<sup>8</sup>

These findings suggest that the false recognition of implied traits reflects spontaneous trait inferences about the person rather than inferences about the behaviors. Other relevant evidence

<sup>7</sup> This analysis could not be performed at the level of participants because (a) the explicit trait judgments were provided by a different group of participants, (b) relating rate of false recognition to RTs for correct rejection is impossible at this level, and (c) each participant only provided a single response per face–trait stimulus.

<sup>8</sup> This last effect in itself is not interesting. It indicates that RTs were slower for trait words for which participants were more likely to make recognition errors. Note also that this effect cannot be interpreted at the level of participants because it relates recognition errors to RTs for correct decisions.

Table 3  
*Effect Sizes ( $r$ ) of Response Times for Correct Rejection of Implied Traits as a Function of Experiment at the Levels of Participants and Stimuli*

| Experiment          | Effect size  |         |
|---------------------|--------------|---------|
|                     | Participants | Stimuli |
| Experiment 2        | .41          | .59     |
| Experiment 3        | .24          | .28     |
| Experiment 4 (10 s) | .23          | .34     |
| Experiment 4 (5 s)  | .27          | .20     |
| Average effect size | .29          | .36     |

comes from the correlation of the false recognition rate with the probability of trait generation given only the behavior. The latter probability was calculated for each behavioral sentence when these sentences were compiled (Uleman, 1988). In contrast to the explicit person judgments, these probabilities are not related to any actor. Participants were asked to read each behavior and to generate a trait, but they were not presented with any actor information or photo. The correlation of the probability with the false recognition of implied traits was practically zero,  $r(24) = -.04$ , whereas the correlation of the latter with the actor judgments was .47,  $p < .025$  (see Table 4).

### Effects of Trait Valence on False Recognition

It is important to address possible effects of trait valence on the false recognition of traits. Trait valence can affect the overall rate of false recognition. For example, previous studies have shown that people are more likely to falsely recognize positive than negative traits (Matlin & Stang, 1978; Todorov, 2002). More important, valence inferences could account for some of the false recognition effects. According to a valence explanation, participants make global, evaluative inferences rather than specific trait inferences. At the recognition test, if the trait word matches the inference evaluatively, participants are likely to falsely recognize the trait. By default, implied traits share the same evaluation as the behavior, and, thus, participants are likely to falsely recognize them. In contrast, if the trait does not match the evaluative inference, participants are less likely to falsely recognize it. To check these two possible effects of valence, we reanalyzed the data from all experiments, coding for the valence of the traits.

Consistent with previous findings, in all six experiments the false recognition of positive traits was higher than the false recognition of negative traits, although the effect was not significant in Experiments 3 and 5. However, this overall valence effect did not qualify the effect of higher false recognition of implied traits relative to randomly paired implied traits (Experiments 1, 2, 5, and 6) or to novel control traits (Experiments 2, 3, and 4).

Another implication of a valence explanation of our false recognition results is that if the control trait (randomly paired implied trait or a novel trait) and the implied trait share the same valence, participants should be equally likely to falsely recognize them. This was not the case. In all six experiments, the evaluative match/mismatch of implied and control traits did not affect the difference between the false recognition of implied traits and the false recognition of control traits. This finding, coupled with the

correlational data above that show that participants made fine distinctions among implied traits, is clearly inconsistent with the valence explanation.

However, the valence explanation could account for the lower rate of false recognition of antonym than control traits. Antonyms are always evaluatively inconsistent with the implied traits. Thus, it is sufficient to retrieve the evaluative inference and to use this inference as a basis of the recognition decision. If this is the mechanism underlying the low rate of false recognition of antonym traits, participants should not discriminate between antonyms and novel control traits that share the same valence. Although the data from Experiment 3 did not support this explanation, the data from Experiment 4, which had more statistical power, were clearly consistent with the valence-matching explanation. The interaction of evaluative match and type of trait was reliable,  $F(1, 185) = 6.29$ ,  $p < .013$ . The difference between the false recognition of evaluatively matched control ( $M = .15$ ) and antonym traits ( $M = .14$ ) was not significant ( $t < 1.00$ ), whereas the difference between evaluatively mismatched control ( $M = .17$ ) and antonym traits ( $M = .11$ ) was significant,  $t(185) = 4.31$ ,  $p < .001$ .

These findings suggest the following mechanism, which posits a role for STIs and spontaneous evaluative judgments, both of which are linked to the actor. Participants do infer the implied trait and link this trait to the actor's representation. However, the actor's face is not a sufficiently strong cue to elicit the spontaneous retrieval of the implied trait. When this trait is presented at the recognition test with the actor's face, participants are likely to falsely recognize it because it matches an existing representation. When a different trait is presented with the actor's face, participants are likely to spontaneously retrieve the evaluation but not the trait itself and, hence, to base their recognition decisions on the retrieved evaluation. Future studies should address the plausibility of this mechanism.

### General Discussion

In six experiments, we showed not only that people unintentionally infer traits from single behaviors but also that they associate these traits with the person who performed the behavior. Experiments 1, 2, 5, and 6 showed that participants were more likely to falsely recognize implied traits when these traits were

Table 4  
*Zero-Order and Partial Correlations Among Explicit Judgments of Behavioral Implications for Traits, False Recognition of Implied Traits, and Response Times for Correct Rejection of Implied Traits*

| Variable                         | Correlation |         |
|----------------------------------|-------------|---------|
|                                  | Zero order  | Partial |
| Judgments/false recognition      | .47**       | .36*    |
| Judgments/response times         | .50**       | .40*    |
| False recognition/response times | .37         | .17     |

*Note.* The partial correlations control for the effect of the third variable on the correlated variables.

\*  $p < .10$ . \*\*  $p < .05$ .

paired with the face of the person who performed the behavior than when these traits were randomly paired with a different familiar face. That was the case even when the number of behaviors was very large (120) and each behavior was presented only for 5 s as well as when participants did not recognize or did not recall the behavior.

Experiments 3 and 4 showed that people go beyond the initial trait inference, in that their person representations were particularly relevant to traits opposite from the implied traits. Participants were less likely to falsely recognize antonyms to implied traits than control traits. Experiment 4 also showed that memory and impression goals engaged different processes, in that doubling the presentation time of stimuli decreased false recognition of implied traits, but only under memory instructions. In all experiments, the analyses at the level of participants and at the level of stimuli were consistent. Across experiments at both the level of participants and the level of stimuli, the effect sizes for false recognition of implied traits were large, and the effect sizes for false recognition of antonym traits were moderate.

The RTs of the correct recognition decisions provided additional support for the hypothesis that participants associated the inferred traits with the person. In Experiments 2, 3, and 4, participants were slower to correctly reject the implied trait than a control trait when it was paired with the actor. The effect sizes for the RTs were moderate.

Finally, correlational analyses showed that the recognition responses reflect trait inferences about the actor. Explicit trait judgments of the actor predicted both the false recognition of implied traits and the RTs for correct rejection of implied traits. That is, the stronger the judgment about the actor was, the stronger the STI was, and the harder (i.e., slower) a decision that required ignoring the trait inference was.

#### *Spontaneous Trait Inferences in the False Recognition Paradigm: Encoding or Retrieval?*

We have argued that the implied traits were inferred during the presentation of the behavior and encoded as part of the person representation. As outlined above, an alternative explanation is in terms of retrieval processes during the recognition test. According to this alternative, participants did not make any inferences at encoding but retrieved the actor's behavior and checked the plausibility of the trait against this behavior in the recognition test. Because randomly paired implied traits and unrelated traits are irrelevant to the behavior, the recognition decision is easy in those cases, and false recognition rates are lower. However, this explanation cannot account for the pattern of RTs findings. The retrieval explanation predicts that RTs for correct rejection of unrelated and randomly paired traits should not differ (they are both irrelevant to the behavior) and that they should be slower than the RTs for implied traits. That was not the case. RTs for rejection of randomly paired implied traits were slower than RTs for unrelated traits and not significantly different from RTs for implied traits. This suggests that implied traits were inferred during encoding. Similarly, the retrieval explanation is inconsistent with the finding that participants were more likely to falsely recognize implied, albeit randomly paired, traits than novel traits (Experiment 2).

Most important, Experiments 5 and 6 provide critical evidence against the retrieval explanation. Even when participants could not

recognize the behavior (Experiment 5) or could not recall it (Experiment 6), they were more likely to falsely recognize implied traits paired with the actor than implied traits randomly paired with different familiar persons. In fact, the effect sizes were practically equivalent for recognized/recalled and unrecognized/unrecalled behaviors, suggesting that the retrieval of behavior did not play any role in the false recognition of implied traits. Both experiments used a very large set of behaviors presented for short time. These are precisely the conditions that are likely to interfere with the successful encoding of the behaviors. It is interesting that comparing the effect sizes for the false recognition of implied traits across experiments (Experiments 1, 2, 3, and 4 vs. 5 and 6) suggests that increasing the number of behaviors does not decrease, at least substantially, the effect of false recognition of implied traits. This is an issue worth directly addressing in future empirical work.

It should be noted that the above findings do not completely rule out all retrieval effects. For example, the difference between the rates of false recognition of implied traits and randomly paired implied traits was higher for recalled behaviors than for nonrecalled behaviors. However, the findings do rule out the possibility that the false recognition of implied traits is entirely or largely due to retrieval effects. The false recognition paradigm constrains participants' responses to two options—accepting or rejecting the trait word. Under these conditions, when participants cannot retrieve the behavior, they have to rely on prior inferences even if these inferences result in only vague familiarity of the trait word.

#### *Spontaneous Antonyms: Going Beyond Initial Inferences*

Two of these studies provide clear evidence not only that people spontaneously infer traits and associate them with those who enact the trait-implying behaviors but also that these inferences clearly imply that the actors lack the opposites of these traits. Two alternative explanations of this effect are that the absence of opposites (antonyms) becomes part of the person representation at encoding or that this absence follows more certainly from the person representation at retrieval and test. The findings from the present studies are more consistent with the latter explanation. First, the effects for implied traits were twice as strong as the effects for antonyms. Second, the RTs for correct rejection of antonyms did not differ from the RTs for unrelated traits, indicating lack of facilitation in the recognition decisions. Third, the correlational analyses showed that explicit person–trait judgments based on the behaviors predicted both the false recognition and the RTs for implied traits but not for antonym traits.

Additional analyses controlling for the evaluative match of the antonym and control traits shed more light on the underlying mechanism. These analyses showed that antonym traits were as likely to be falsely recognized as were novel unrelated traits if they shared the same valence. This finding suggests that the presentation of the actor's face was sufficient to trigger the retrieval of the evaluation but not sufficient to trigger the retrieval of the specific trait and that the recognition decisions were based on the retrieved evaluation.

It is interesting that this mechanism is inconsistent with trait rating results reported by Skowronski et al. (1998) and Mae et al. (1999). Their research on spontaneous trait transference suggests that implications of spontaneously inferred traits probably do not include general halo effects. Inferred traits' valence alone did not

affect other ratings. However, it is possible that these differences are related to the nature of the tasks. Rating tasks are explicit judgment tasks in which participants can be guided by naive theories of what constitutes an unbiased judgment (e.g., Wegener & Petty, 1997). Perhaps word recognition tasks are less susceptible to such influences and are more likely to reveal valence effects.

### *Spontaneous Trait Inferences or Spontaneous Behavioral Inferences?*

The findings of the six experiments show that participants engaged in spontaneous inferences about the actor. However, in both the false recognition and the savings paradigms there is some unresolved ambiguity about the nature of these inferences. For example, one may argue that participants made inferences about the behavior and that these inferences about the behavior (rather than the actor) were linked to the actor. That is, the inference is that the actor performed an honest act rather than that the actor is honest. The strongest evidence against this assumption comes from the correlational analyses at the level of the stimuli. These analyses showed that explicit trait judgments of the actor predicted STIs. In contrast, measures of the trait implications of the behaviors, which did not include the actors, did not predict these inferences. These findings do suggest that participants made trait inferences about the actor rather than merely about the behavior. However, additional studies and other ways of distinguishing between these two kinds of inferences are needed.

### *Representing Actor-Linked Spontaneous Trait Inferences*

Presenting actors' faces paired with implied traits clearly resulted in a high rate of false recognition of these traits. Does this finding suggest that the trait would be spontaneously retrieved if only the actor's face were presented? The post hoc analyses controlling for the evaluative matching of the probe trait with the implied trait suggested that such spontaneous retrieval without additional prompts is unlikely. People are likely to retrieve the evaluation of the implied trait but not the trait itself. If this account is correct, the implications of spontaneously inferred traits could be rather limited compared with the implications of explicitly inferred traits. On the other hand, if STIs are only implicitly represented in the sense of not being spontaneously retrievable, their effects would not be easily detected and could result in a variety of person judgment biases. These hypotheses could be tested in experiments directly comparing the effects of spontaneously inferred traits and explicitly inferred traits using different explicit and implicit judgment tasks.

### *The Spontaneous Binding of Trait Inferences to Persons' Faces*

The six experiments using the false recognition paradigm provide convergent evidence with the savings paradigm (Carlston & Skowronski, 1994; Carlston et al., 1995) that STIs are bound to actors' faces in long-term memory. We stressed the differences between the savings and the false recognition paradigms in the introduction of this article. However, these paradigms share features that suggest why the observed effects of associating STIs and actors are so robust. In contrast to paradigms such as cued recall,

both these paradigms use faces as critical stimuli, and both paradigms exploit actor-to-trait links as well as trait-to-actor links in the experimental tasks.

There may be something uniquely important about faces as person representations. Among other primates—where language and naming play no role—success in navigating the complex social structure of kin relations, dominance hierarchies, and alliances depends fundamentally on keeping individual identities straight, and all primates do this very well (e.g., Tomasello & Call, 1997, especially Chapter 7). The same is certainly true of humans (M. H. Johnson & Morton, 1991). Facial features are distinctive (Farah, Wilson, Drain, & Tanaka, 1998), and our species (and other primates) has a long history of relying on this distinctiveness to organize social knowledge. Consistent with the functional importance of facial recognition, there are neurological structures devoted specifically to face recognition. Damage to them results in prosopagnosia (e.g., Farah, Rabinowitz, Quinn, & Liu, 2000). All of this suggests a particular readiness to organize person information around facial representations of others.

Second, unlike most previous research on binding STIs, the procedures in the false recognition and savings paradigms used persons (faces) rather than traits as cues or prompts. If faces are uniquely central to the organization of person information, then person–trait bonds may be asymmetric, stronger from person (face) to trait than vice versa. A person's face can cue the retrieval of multiple traits possessed by the person. Moreover, retrieved traits that are part of the person representation can facilitate, in turn, the retrieval of other person traits (Srull, 1981). Thus, thinking of a person may bring specific traits to mind more readily than thinking of a trait brings specific persons to mind. These studies were not designed to detect such an asymmetry, and they provide no direct evidence of it. But this merits future research.

### *Conclusion*

In this article, we introduced a false recognition paradigm to study STIs. This paradigm produces reliable and large effects showing that STIs refer to the actor. People not only spontaneously infer traits but also uniquely associate these inferences with the actor. Such inferences are individuated and encoded into the actor's representation and serve as a basis for other person-related cognitions.

### *References*

- Carlston, D. E., & Skowronski, J. J. (1994). Saving in the relearning of trait information as evidence for spontaneous inference generation. *Journal of Personality and Social Psychology*, 66, 840–856.
- Carlston, D. E., Skowronski, J. J., & Sparks, C. (1995). Savings in relearning: II. On the formation of behavior-based trait associations and inferences. *Journal of Personality and Social Psychology*, 69, 420–436.
- D'Agostino, P. R. (1991). Spontaneous trait inferences: Effects of recognition instructions and subliminal priming on recognition performance. *Personality and Social Psychology Bulletin*, 17, 70–77.
- Duff, K. J., & Newman, L. S. (1997). Individual differences in the spontaneous construal of behavior: Idiocentrism and the automatization of the trait inference process. *Social Cognition*, 15, 217–241.
- Farah, M. J., Rabinowitz, C., Quinn, G. E., & Liu, G. T. (2000). Early commitment of neural substrates for face recognition. *Cognitive Neuropsychology*, 17, 117–123.

- Farah, M. J., Wilson, K. D., Drain, M., & Tanaka, J. N. (1998). What is "special" about face perception? *Psychological Review*, 105, 482–498.
- Gross, D., Fischer, U., & Miller, G. A. (1989). The organization of adjectival meanings. *Journal of Memory and Language*, 28, 92–106.
- Hamilton, D. L., Katz, L. B., & Leirer, V. (1980). Organizational processes in impression formation. In R. Hastie, T. M. Ostrom, E. B. Ebbesen, R. S. Wyer, D. L. Hamilton, & D. E. Carlston (Eds.), *Person memory: The cognitive basis of social perception* (pp. 121–153). Hillsdale, NJ: Erlbaum.
- Johnson, M. H., & Morton, J. (1991). *Biology and cognitive development: The case of face recognition*. Oxford, England: Blackwell.
- Johnson, M. K., Hashtroudi, S., & Lindsay, D. S. (1993). Source monitoring. *Psychological Bulletin*, 114, 3–28.
- Johnston, R. A., & Bruce, V. (1990). Lost properties? Retrieval differences between name codes and semantic codes for familiar people. *Psychological Research*, 52, 62–67.
- Luce, R. D. (1986). *Response times: Their role in inferring elementary mental organization*. Oxford, England: Oxford University Press.
- Mae, L., Carlston, D. E., & Skowronski, J. J. (1999). Spontaneous trait transference to familiar communicators: Is a little knowledge a dangerous thing? *Journal of Personality and Social Psychology*, 77, 233–246.
- Matlin, M. W., & Stang, D. J. (1978). *The Pollyanna principle: Selectivity in language, memory, and thought*. Cambridge, MA: Schenkman.
- McWeeny, K. H., Young, A. W., Hay, D. C., & Ellis, A. W. (1987). Putting names to faces. *British Journal of Psychology*, 78, 143–149.
- Park, B. (1989). Trait attributes as on-line organizers in person impressions. In J. N. Bassili (Ed.), *On-line cognition in person perception* (pp. 39–59). Hillsdale, NJ: Erlbaum.
- Reeder, G. D., & Brewer, M. B. (1979). A schematic model of dispositional attribution in interpersonal perception. *Psychological Review*, 86, 61–79.
- Rosenberg, S., & Sedlak, A. (1972). Structural representations of implicit personality theory. In L. Berkowitz (Ed.), *Advances in experimental social psychology* (Vol. 6, pp. 235–297). New York: Academic Press.
- Schneider, D. J. (1973). Implicit personality theory: A review. *Psychological Bulletin*, 79, 294–309.
- Skowronski, J. J., Carlston, D. E., Mae, L., & Crawford, M. T. (1998). Spontaneous trait transference: Communicators take on the qualities they describe in others. *Journal of Personality and Social Psychology*, 74, 837–848.
- Skowronski, J. J., & Shook, J. (1997). Facilitation in repeated trait judgments: Implications for the structure of trait concepts. *Journal of Experimental Social Psychology*, 33, 21–46.
- Srull, T. K. (1981). Person memory: Some tests of associative storage and retrieval models. *Journal of Experimental Psychology: Human Learning and Memory*, 7, 440–463.
- Todorov, A. (2002). Communication effects on memory and judgments. *European Journal of Social Psychology*, 32, 531–546.
- Tomasello, M., & Call, J. (1997). *Primate cognition*. New York: Oxford University Press.
- Tulving, E., & Thomson, D. M. (1973). Encoding specificity and retrieval processes in episodic memory. *Psychological Review*, 80, 352–373.
- Uleman, J. S. (1988). [Trait and gist inference norms for over 300 potential trait-implying sentences]. Unpublished raw data.
- Uleman, J. S., & Moskowitz, G. B. (1994). Unintended effects of goals on unintended inferences. *Journal of Personality and Social Psychology*, 66, 490–501.
- Uleman, J. S., Newman, L. S., & Moskowitz, G. B. (1996). People as flexible interpreters: Evidence and issues from spontaneous trait inference. In M. P. Zanna (Ed.), *Advances in experimental social psychology* (Vol. 28, pp. 211–279). San Diego, CA: Academic Press.
- Van Overwalle, F., Drenth, T., & Marsman, G. (1999). Spontaneous trait inferences: Are they linked to the actor or to the action? *Personality and Social Psychology Bulletin*, 25, 450–462.
- Wegener, D. T., & Petty, R. E. (1997). The flexible correction model: The role of naive theories of bias in bias correction. In M. Zanna (Ed.), *Advances in experimental social psychology* (Vol. 29, pp. 141–208). New York: Academic Press.
- Winter, L., & Uleman, J. S. (1984). When are social judgments made? Evidence for the spontaneousness of trait inferences. *Journal of Personality and Social Psychology*, 47, 237–252. Also see correction in *Journal of Personality and Social Psychology* (1986), 50, 355.
- Zarate, M. A., Uleman, J. S., & Voils, C. I. (2001). Effects of culture and processing goals on the activation and binding of trait concepts. *Social Cognition*, 19, 295–323.

Received March 12, 2001

Revision received February 21, 2002

Accepted April 10, 2002 ■

## Wanted: Your Old Issues!

As APA continues its efforts to digitize journal issues for the PsycARTICLES database, we are finding that older issues are increasingly unavailable in our inventory. We are turning to our long-time subscribers for assistance. If you would like to donate any back issues toward this effort (preceding 1982), please get in touch with us at [journals@apa.org](mailto:journals@apa.org) and specify the journal titles, volumes, and issue numbers that you would like us to take off your hands.