

Self-Serving Interpretations of Flattery: Why Ingratiation Works

Roos Vonk
University of Nijmegen

Persons who are flattered are more likely to assign credibility to and like the flatterer than observers, presumably because they are motivated by vanity. In existing studies, however, the difference between targets and observers has been confounded with other variables. The present experiments demonstrate that the target–observer difference in judgments of an ingratiation is not affected by these confounding variables, such as cognitive resources, the motive to like one's interaction partner, or to form an accurate impression, or mood. Results further suggest that, whereas cognitive responses to ingratiation are different among participants with high versus low self-esteem, affective responses and judgments of the ingratiation are not qualified by any personality variables.

Vanity, a college professor, meets a new student, Sly, who has been assigned to her as his supervisor. During their first meeting, Sly says he is very excited about working with Vanity because he admires her work a great deal and her classes are the best. He enthusiastically endorses all her suggestions about his thesis and expresses great happiness that she is his supervisor.

Vanity can draw two conclusions from Sly's behavior. One is a correspondent inference: Sly has quite a favorable impression of her. As a result, Vanity will probably like him as well (cf. Backman & Secord, 1959). Alternatively, she may be suspicious of Sly's motives: Possibly, he is flattering her because he depends on her (e.g., for his grade, future perspectives, recommendations). In this case, her inference will be more moderate because she is uncertain if his behavior is guided by ulterior motivation (Fein, Hilton, & Miller, 1990). Put differently, she engages in situational correction for her initially favorable inference by taking into account her power over Sly as a situational cue for his behavior (cf. Vonk, 1998b). This correction requires more cognitive effort than a correspondent inference (Gilbert, Pelham, & Krull, 1988). Thus, she would require both the opportunity (e.g., processing time) as well as the motivation to engage in this process.

As for motivation, Vanity is probably willing to expend some effort in forming an accurate impression of Sly because she is going to be working with him. On the other hand, her cognitive resources are restrained by the fact that she is not a passive observer but an active participant in the conversation with Sly (cf. Gilbert et al., 1988). Thus, she is expending cognitive resources on managing her part of the interaction. Moreover, Sly is flattering her, so it is tempting to take his praise at face value. If she does

this, everybody is happy: Sly has managed to make a friendly, positive start in the collaboration with his supervisor; Vanity's ego is bolstered; Sly likes her, and she likes him, which is good because they are going to have to spend some time together.

Generally, when people expect further interaction with someone, they are motivated to form an accurate impression of this person. As a result, they engage in more attentive and elaborate information processing (e.g., Berscheid, Graziano, Monson, & Dermer, 1976; cf. Erber & Fiske, 1984; Vonk, 1998a). However, in the example above, at least three variables compete against accuracy: (a) limited cognitive resources, (b) the motive to like the interaction partner, and (c) the motive to be liked and admired. All of these variables may reduce the likelihood that the sincerity of Sly's behavior is questioned and that a thoughtful attributional analysis is instigated.

In light of these variables, we should not be surprised to learn that people generally tend to like those who flatter them. In a meta-analysis, Gordon (1996) found that people form more favorable impressions of an ingratiation when the ingratiation (e.g., opinion conformity, flattery) is directed toward them than when they are uninvolved observers. This result appears to corroborate Jones's (1964) seminal notion of "vain distortion" (p. 77): It is more rewarding to assign credibility to lavish praise directed toward the self than toward someone else. As a consequence, "the targets of ingratiation . . . may be less sensitive to the implications of ulterior motivation than bystanders" (Jones, Stires, Shaver, & Harris, 1968, p. 350). To be precise, the effect demonstrated pertains not to the credibility of the flatterer, but to liking. We may assume, however, that these variables are highly correlated. A person who is seen as sincere tends to be seen as likeable as well. Moreover, in case of ingratiation, credibility implies that the flatterer truly likes or admires the target, and the target is likely to reciprocate those feelings.

In the extant literature, it is assumed that the target–observer effect above is caused by the target's motivation to be flattered: When people are the target of ingratiation, their self-esteem is served by accepting the flattery uncritically; on the other hand, when they are observers, their ego is not at stake and they may examine the ingratiation's behavior more critically. For instance, Jones (1990, p. 179) noted: "The best evidence that vanity plays a

Experiments 1–3 were conducted while Roos Vonk was working at Leiden University, Leiden, the Netherlands. Thanks to Inge Feenstra, Fleurike Kamerbeek, Claudia Schoutens, and Annemiek van der Weel for assisting in these experiments; and to Glenn Reeder, Constantine Sedikides, Tim Wilson, and three anonymous reviewers for helpful comments on this article.

Correspondence concerning this article should be addressed to Roos Vonk, Department of Social Psychology, University of Nijmegen, P.O. Box 9104, 6500 HE Nijmegen, the Netherlands. E-mail: vonk@psych.kun.nl

role in episodes of ingratiation comes from experiments comparing the reaction of . . . targets . . . with that of . . . observers or bystanders." From this perspective, the target–observer effect may be seen as a specific variant of a more general phenomenon: People tend to instantly accept statements that converge with their desired beliefs, whereas information that contradicts such beliefs is examined more stringently (e.g., Ditto & Lopez, 1992; Kunda, 1990; Pyszczynski & Greenberg, 1987). Considering that people generally prefer to have a positive self-view, this means they are willing to readily accept positive statements about the self, without giving much thought to the motives of the person who makes the compliment.

However, as noted above, several other variables are involved in the interaction between ingratiation and target, and many of these may account for the target–observer effect as well. Contrary to Jones's (1990) conclusion, differences between targets and observers do not provide any evidence at all that vanity plays a role in interpreting ingratiation behavior. In fact, if we take a critical look at the studies that have demonstrated the effect, it may be noted that the variable of interest (i.e., whether the participant is the target of flattery or is an uninvolved observer) is seriously confounded with one or more other variables in all studies.¹

First, in most studies where participants were the target of ingratiation, participants were also cognitively busy, because they were sending messages to alleged other participants (e.g., Fodor & Farrow, 1979; Jones, Gergen, & Jones, 1963; Kipnis & Vanderveer, 1971) or were involved in an interaction with the ingratiation (e.g., Baron, 1986; Godfrey, Jones, & Lord, 1986; Jones et al., 1968, Study 1). Because of the resultant cognitive limitations, it is possible that these participants were unable (or at least less able than their uninvolved counterparts in control conditions) to engage in situational correction, or in the sophisticated attributional analysis that is required to consider ulterior motives (Fein, 1996), even if they had been motivated to.

Second, as noted earlier, a robust finding is that people are motivated to like persons they interact with or expect to interact with later (e.g., Berscheid et al., 1976; Darley & Berscheid, 1967; Tyler & Sears, 1977; see Vonk, 1998a). Thus, (expected) interaction produces increased liking even if the interaction partner is not engaging in ingratiation at all. In ingratiation studies, this means that target participants may have been motivated to like the ingratiation simply because they were interacting with this person or were expecting an interaction, whereas observers were not.

Third, it is possible that being flattered simply puts targets in a good mood even before they start questioning the flatterer's motives. This positive mood may have two consequences. First, they may scrutinize information less critically (cf. Bless, Mackie, & Schwarz, 1992; Bodenhausen, 1993), so that the correspondence bias is increased (Forgas, 1998). Second, subsequent judgments of virtually any stimulus, including the flatterer, may be more positive (cf. Forgas & Bower, 1987).

Finally, because psychologically healthy people have high self-esteem, it is important to realize that the positive feedback given by an ingratiation will usually be evaluatively consistent with the self-concept of the target. And, just as consistent information about others is studied less carefully than inconsistent information (e.g., Belmore, 1987; Vonk, 1994), so may consistent information about the self be accepted at face value, without giving it much thought, simply because it confirms something one already knew anyway.

Uninvolved observers, on the other hand, do not have an existing expectancy of the person being ingratiation, so the flattery does not match anything that was already known. According to this view, the target–observer effect is caused mainly by cognitive and not motivational factors: People simply do not question or elaborate on information that fits with their existing beliefs.

The purpose of the present experiments is twofold. First, it will be examined whether the target–observer effect still occurs when there are no differences between targets and observers in information processing capacity and accuracy motivation. Therefore, in all experiments reported in this article, participants did not interact with others in any way, so that their attention was not absorbed by the interaction, and both targets and observers were given ample time to observe and think about the ingratiation episode. Regarding accuracy motivation, this variable is manipulated in Experiments 1 and 2. Second, assuming that the effect still occurs in these conditions, some of the possible underlying processes will be examined. In Experiment 3, the role of mood is considered; in Experiment 4, the evaluative consistency between the ingratiation comments and preexisting knowledge of the target is under consideration.

All experiments included a manipulation of *self-relevance*: Participants read a flattering description by an alleged other participant, the ingratiation, that was either about themselves (high self-relevance) or someone else (control: no self-relevance). Participants were made aware that the ingratiation had a motive to flatter the person described because she or he depended on this person. Also, participants were aware that the ingratiation knew the description would be read by the person described. Thus, the favorable description could be attributed either to genuine liking and admiration by the ingratiation, or to ulterior motivation. The main dependent variables were participants' liking for the ingratiation, whether they perceived the ingratiation as "slimy,"² and their assessments of whether the ingratiation's description was sincere.

¹ In addition, the meta-analysis by Gordon (1996) includes some studies that are not relevant to the issue under consideration here (i.e., the dilemma of accepting or questioning flattery), because it was clear that the ingratiation was sincere. For instance, in several studies the actor was not dependent on the target (e.g., Byrne & Rhamey, 1965) or clearly expressed his or her true beliefs (e.g., Tjosvold, 1978, who deliberately took measures to make participants assume that the ingratiation was unaware that they would see his judgments of them), so that participants did not have any reason to suspect the ingratiation of ulterior motivation.

² "In the Netherlands, where these studies were conducted, there is no general word for ingratiation. The closest resemblance is derived from the noun 'slime' [slijm], which literally refers to the same slippery substance as the English word, but is more often used to describe persons and behaviors. The verb 'to slime' [slijmen] refers to the behavior of ingratiation oneself for ulterior motives. . . . As in English, a person who engages in this type of behavior is described by the adjective 'slimy' [slijmerig]. The word 'slime' and its conjugations have a negative connotation and are used frequently to describe flattery, overly friendly behavior, and brown-nosing. So, for all practical purposes, these words refer to the same class of behaviors as the term ingratiation, but they are much more informal" (quoted from Vonk, 1998b, pp. 849–850).

Experiment 1

Experiments 1 and 2 focus on the role of motivational variables in interpreting ingratiating behavior. As noted above, the fact that people are usually interacting with an ingratiation may affect their interpretations about the behavior, because (expected) interaction induces several motives, such as the motive to form an accurate impression as well as the motive to like the interaction partner. To examine the possible effects of these motives, participants in Experiment 1 either did or did not expect to interact with the ingratiating actor. For no-interaction conditions, the prediction was that the target–observer effect would occur, that is, participants who were themselves the targets of flattery would form more favorable impressions of the actor than those who observed that someone else was being flattered (Hypothesis 1).

Regarding the comparison with conditions where participants did expect to interact with the actor, two competing hypotheses were formulated. Given sufficient processing capacity, as in these studies, it is possible that blatant ingratiation is identified more easily when perceivers expect further contact with the actor: The anticipated interaction increases the motivation to form an accurate impression; this makes it more likely that they engage in effortful processing (e.g., Neuberg & Fiske, 1987) and perform a situational correction (Vonk, 1999b). This should occur even when participants are themselves the target of ingratiation, because they are motivated to avoid being fooled by their prospective interaction partner. As a result of this accuracy orientation, the target–observer effect should be reduced in interaction conditions. Further, compared with no-interaction conditions, the ingratiation should be judged as less likeable and sincere, and as more slimy, because participants expecting an interaction are more suspicious of the ingratiation's motives (Hypothesis 2A).

Alternatively, perceivers are motivated to form a favorable impression of their interaction partner. This motive would be foiled by assuming that they are being ingratiated (or that someone else is, in control conditions), because ingratiation toward people that one depends on is seen as socially undesirable and insincere (cf. Vonk, 1999a, 2001). Thus, regardless whether the actor's flattery is directed toward themselves or someone else, participants expecting an interaction may be motivated to infer that the description is sincere. Compared with no-interaction conditions, this would produce more favorable judgments of the actor (Hypothesis 2B). Note that both Hypotheses 2A and 2B predict a reduction of the target–observer difference in interaction conditions, but according to Hypothesis 2A, judgments in these conditions will be less favorable than in no-interaction conditions, whereas Hypothesis 2B predicts they are more favorable.

Method

Participants and Design

The design was a 2 (self-relevance) $\times$ 2 (interaction expectancy) factorial design. Participants were 80 undergraduates with different majors (58 women, 22 men). They were paid for participating in this study and in an unrelated scenario study on social justice. They participated in groups of 6 to 8 at a time. In some cases, they saw some of the other participants shortly before the experiment, but they did not know each other.

Introduction of the Study

Participants were recruited during classes to participate in a series of studies about different topics, including a study on "how people form impressions of others and how they collaborate with others when working at a task." On arrival, the experimenter asked each participant to read and sign a "confidentiality form." This was done to make it credible that participants would later receive information about each other. The form explained that participants would fill out a personality test, and that in some cases, the results of their test would be presented to another participant; this would always be a participant whom they did not know because this participant had a different major. They were asked to sign a statement declaring (a) they consented that their test results be shown to another participant, and (b) they would not discuss any information they received about other participants with anyone. All participants complied with this request.

The experimenter then seated the participants in an individual cubicle with a computer that paced them throughout the entire experiment. By pressing the *Return* key, the participant determined the pace of all text on the screen. The experiment started by asking participants to type in their first name and gender. The instructions explained that their first names were used only during the experiment, so that we would not have to denote everyone by means of numbers, and that the names were retained only in the computer's working memory and would not be saved.

The study consisted of several parts. The first part was a series of items from different personality tests. Participants were reminded that their responses would be treated confidentially and were urged to answer the questions as truthfully as possible. They were told that everyone who participated was taking this test. The primary goal of this first part of the study was to convey to participants that the flattering description they would later read was based on some personality information. The personality test consisted of 10 questions assessing self-esteem (constituting a Dutch adaptation of Rosenberg's, 1965, self-esteem scale; cf. Vonk & Ashmore, 1993), 10 questions on interpersonal orientation, 5 questions from a test on Machiavellism, and 8 questions from a test of personal values. In actuality, only the responses on the self-esteem items were registered.

After completing these questions, participants learned that there would be two groups. Participants in one group would see the test responses given by another participant (which would be transferred to them by the computer server to which they were connected) and would be asked to write a description of their impression of this person. Participants in the other group would get to read one of these descriptions. In actuality, the former group did not exist. All participants were told that they had been randomly assigned to the latter group, and that they would participate in a different and unrelated study while the participants in the other group looked at the test results of others and wrote a description of their impressions. Subsequently, the study on social justice was conducted, which took about 25 min.

Manipulations and Cover Story

Self-relevance. When the experiment was resumed, the instructions on the screen reminded participants that, in the meantime, other participants had looked at someone's test responses and had typed in a description of their impression of that person. Participants would read one of these descriptions. Subsequently, they were informed of the name of the participant who had written the description they would read (the ingratiation), and the name of the participant whom the description was about (the target). The name of the ingratiation was always Ronald or Laura, depending on the participant's gender: All participants read about an actor of their own sex. In the self-relevance condition, the name of the target was the participant's own name. Thus, the text on the screen specified: "You, [participant's name], are going to read a description that Laura (Ronald) has written about [participant's name], that is, about yourself." In the control condi-

tion, the text read: "You, [participant's name], are going to read a description that Laura (Ronald) has written about Francis (Frank), that is, about another participant whom you do not know."

Dependence of the ingratiation. Participants then received more information about the circumstances in which the description by the ingratiation had been written. Because this cover story was rather complex, they were urged to read this information carefully; they were allowed to take notes, and they were provided with several opportunities to reread parts of the information. In all conditions, it was explained that we were interested in how people describe others on whom they depend. Therefore, they were told, some participants in the study had been misled to believe that they would work on a task with this person, and that this person was going to determine how much money they would receive for the task. They had also been told that this person would read their description of him or her before working on the task.

In this way, participants were led to believe that the ingratiation (a) expected to be dependent on the target and (b) expected the target to read his or her impression description. This "double bluff" (i.e., participants were misled to believe that the actor had been misled) was introduced because it avoids a confounding between self-relevance and expected interaction: If participants were to assume that the ingratiation was indeed going to work on a task with the target of the description, it would imply that participants who read a description about themselves expected to work with the target, whereas those in the control condition did not; in this case, expected interaction and self-relevance would not be orthogonal.

Expected interaction. In no-interaction conditions, it was stressed that the story presented to the ingratiation, about the alleged upcoming task, was not true: In actuality, there would not be an additional task, nobody would decide how much money someone else would receive, and all participants would receive the same financial compensation.

In expected-interaction conditions, it was explained that the story was only partially true. Participants who were going to read a description about themselves (self-relevance conditions) were told that they would be working on a task with the ingratiation, but they would *not* determine the ingratiation's financial compensation, and everyone would be rewarded similarly, regardless of their performance at the task.

Participants in the control condition were told that everyone would be paired with another participant to work on a task, but not in the way it had been described to the ingratiation. In actuality, participants in the other group would work with the person whose impression description they had read. Thus, Laura (Ronald) would be working with the participant, not with Frances (Frank). Here, too, participants were told that nobody would determine anyone else's financial compensation.

After this explanation, the entire cover story was summarized. The summary first repeated what was actually going to happen, and then what the ingratiation had been led to believe. This summary was concluded by saying, "The reason we are telling you all this, is that the description you will read might contain references to things that are not really going to happen. For instance, Laura (Ronald) might mention something about the money for the task. If you are not informed of what we told her (him) in advance, you would not understand any of such remarks." (In actuality, the ingratiation did not mention the money at all, but this way participants were reminded of the actor's dependence and were simultaneously given a plausible reason for why we let them in on how the ingratiation had been misled.)

Finally, participants were informed of the specific instructions that had been given to the ingratiation. For instance, in the self-relevance conditions, they were told, "First, Laura (Ronald) has been told that she will later work with you and that you will determine how much money she receives for this study. Second, we have explained to Laura the different parts of the personality test and have shown her your responses on each part. Finally, she has been asked to type in her impression of you and her expectancies about working with you. We also told her that you were going to read this description before the two of you would do the task." Subsequently, the

computer ostensibly started searching for files containing an impression description, and stopped the search as soon as the ingratiation's data had been found.

Target Description

The ingratiation's impression description was the same in all conditions and contained two components of ingratiation: other-enhancement and expression of agreement and similarity. To appear realistic, the description contained three typing errors. The text read (translated) as follows:

I am Laura (Ronald) and I understand that I will work with . . . [participant's name or other name]. I have seen her test results. I noticed that we have the same ideas about many things. Especially on how we deal with other people we are very much alike. It seems to me that anyone can get along with [name], but I will certainly get along great with her myself. I'm very glad that I will work with [name] because we seem to be so similar. And I'm sure that she will take her responsibility very seriously and that she'll do an excellent job at whatever it is we're doing, because she seems to have many talents. Also, she seems like a great person to go out with or to talk with, but anyway, that's not the issue now. I think she'll find the right way to do the task. She's given many responses that appeal to me. All in all, my impression is that she's a really nice person, easy to get along with, and someone who has many qualities. I'm looking forward to meeting her.

Dependent Variables

After reading the description, participants were asked to indicate their liking for the ingratiation (1 = *dislikeable*, 7 = *likeable*) and to rate the ingratiation on a series of 7-point trait scales, including *slimy*. They also indicated to what extent the ingratiation's description was sincere. The instructions stressed that all questions were solely for the investigators and that participants' responses were strictly anonymous and would not be shown to any other participant. Several manipulation checks were added; for example, participants had to indicate who would be working with whom, what their hierarchical relationship was, and who had given an impression description of whom.

Finally, to probe for suspicion, participants were asked to type in their ideas about the goal of the study. After this, they were paid and debriefed.

Results and Discussion

One participant was discarded because she had accidentally typed in the wrong gender code, so that she got to read about a male ingratiation and all the text referred to herself as a male. One participant, who had been assigned to the high self-relevance condition, was dropped because he expressed suspicion in his response to the last, open-ended question, by doubting that Ronald really existed. None of the other participants appeared suspicious; the large majority assumed, as they were told, that the goal of the study was to examine whether being dependent on a person affects one's description of that person. In expected-interaction conditions, some assumed that we were interested in the question of how prior information about a person affects the way in which people work together on a task.

Participants' gender did not affect the results—there was only a marginal main effect of gender, $F(3, 68) = 2.56, p < .07$, such that judgments of female subjects about female targets were more positive than judgments of males—and was dropped from the analyses reported below.

In this experiment, as well as all subsequent experiments, self-esteem ($\alpha = .80$ across all four experiments) was dichotomized and included in the analyses of variance (ANOVAs). However, in analyses on judgments of the ingratiation, it produced a significant effect on only one dependent variable in only one of the studies. Therefore, this variable was dropped from the analyses reported here.

A 2 (expected interaction) $\times$ 2 (self-relevance) multivariate analysis of variance (MANOVA) on the likeability ratings (1 = *dislikeable*, 7 = *likeable*), slime ratings (1 = *not at all*, 7 = *highly*), and ratings of whether the description was sincere (1 = *not at all*, 7 = *highly*) produced a main effect of expected interaction, $F(3, 72) = 5.98, p < .01, \eta^2 = .20$; and of self-relevance, $F(3, 72) = 2.92, p < .05, \eta^2 = .11$. The two-way effect was nonsignificant ($F < 1$). Compared with no-interaction conditions, participants who expected to interact with the ingratiation judged the ingratiation as substantially more likeable (5.82 vs. 4.47), $F(1, 74) = 15.90, p < .01$; less slimy (4.09 vs. 4.96), $F(1, 74) = 6.83, p < .05$; and more sincere (5.18 vs. 3.82), $F(1, 74) = 17.44, p < .01$. Compared with observer conditions, participants who were the target of ingratiation judged the ingratiation as more likeable (5.23 vs. 4.84), $F(1, 74) = 1.58, ns$;³ less slimy (4.20 vs. 5.00), $F(1, 74) = 6.11, p < .05$; and more sincere (4.75 vs. 4.03), $F(1, 74) = 4.65, p < .05$.

The present results replicate the finding that targets of ingratiation form more favorable judgments of the ingratiation than observers. In particular, it appears that targets were less likely than observers to question the sincerity of the ingratiation's statements, and less likely to infer that the person who depended on them was being slimy, that is, was ingratiating them. This effect emerged even though participants in this study did not have to allocate their cognitive resources to any other task in addition to observing the ingratiation's behavior, and they had the opportunity to think about the ingratiation's description for as long they wanted to. Thus, participants were able to use the cognitive resources required to engage in the thoughtful attributional analysis that is inherent to situational correction (Gilbert et al., 1988) and suspicion of ulterior motivation (Fein, 1996).

We may tentatively conclude, then, that processing capacity is not a crucial variable in the target–observer effect. Regarding interaction motives, a similar conclusion can be drawn. Participants who expected to interact with the ingratiation were generally more positive about this person, indicating that expected interaction enhances liking (e.g., Berscheid et al., 1976; Vonk, 1998a), but this effect was not qualified by self-relevance. Thus, the effects of expected interaction and self-relevance are additive. Although this could lead to an inflated target–observer effect in studies where these two variables are confounded, the present data suggest that self-relevance is sufficient to produce the effect.

Experiment 2

In Hypothesis 2, a distinction was made between two competing possibilities: Because expected interaction could have the effect of increasing the motive to like the ingratiation, as well as increasing accuracy motivation, it was conceivable that judgments in interaction conditions could be either more or less favorable than in control conditions. The results of Experiment 1 show that the motive to like the interaction partner was clearly predominant.

Thus, it seems that participants in interaction conditions were far more motivated to arrive at a favorable impression than an accurate impression, possibly because they did not depend on their interaction partner anyway (cf. Vonk, 1998a). Thus, even though the results of the study rule out the possibility that the expected interaction with a dependent ingratiation accounts for the target–observer effect, they do not reveal whether accuracy motivation can reduce the effect. In everyday life, even people who are in power may have something to lose when they uncritically embrace an ingratiation's flattery. For instance, rich people stand to lose money by dating or marrying an ingratiation who is more interested in their wealth than any other quality; supervisors stand to lose their position by trusting a young and ambitious subordinate who is after their job (cf. Fein & Hilton, 1994). In these cases, it is possible that people are more careful, and less likely to be guided by the motive to like the person they interact with. Having examined the role of the motive to like in Experiment 1, Experiment 2 was designed to study the role of accuracy motivation. Specifically, the experiment examined whether the target–observer effect is reduced when participants stand to lose something by trusting the ingratiation.

Method

Participants and Design

A 2 (self-relevance) $\times$ 2 (stakes: high, low) factorial design was used. Participants were 125 undergraduates with different majors (82 women, 42 men), recruited during classes. They were paid 15 guilders (at the time, about \$7.50) for participating in this study and an unrelated filler study. One participant was discarded because he had already taken part in Experiment 1.

Procedure

Participants were recruited to participate in a study on how people form impressions of others. The procedure was similar to Experiment 1. There were a few minor changes in the personality test administered in the first part of the study. The most important difference with Experiment 1 concerns the expected relationship between the ingratiation and the target. In all conditions, it was clear that the ingratiation and the target would never meet. Subjects were told that some of the other participants, namely, those who had been assigned to write an impression description on the basis of someone's test results, had been recruited to participate in this study for 5 guilders (i.e., about \$2.50), instead of 15. Thus, Laura/Ronald (the ingratiation) was only getting 5 guilders, whereas the subject was getting 15. These other participants had already agreed to participate for 5 guilders, so, subjects were told they need not feel concerned about it. However, at the end of the experiment, the person described by the ingratiation (either the subject or Frances/Frank, depending on self-relevance) would be asked if she or he was willing to give 5 extra guilders to Laura/Ronald (the ingratiation). The ingratiation knew that this request would be made. As in Experiment 1, subjects were told that we had done this because we were interested in the effects of dependence on the impression descriptions.

Hence, in conditions where subjects were themselves the target of ingratiation, they would be the ones deciding whether the ingratiation would receive 10 guilders instead of 5. They were told that they had no obligation

³ In this particular study, the effect of self-relevance on likeability was not significant. However, the effect was significant in all other experiments, including ones that are not reported here.

whatsoever to do this, because the other participant had already agreed to work for 5 guilders; also, they would never see this person, so they need not be concerned about his or her response. However, they were told, after the experiment they might be willing to assign 5 guilders to this person “simply because you feel like it, or because it seems more fair, or because you like Ronald/Laura.”

In high-stakes conditions, subjects were told that, if they decided to give 5 guilders to the ingratiation, they would have to give up 5 guilders themselves, so that both of them would end up with 10 guilders. (“After all, if you are the director of your own company, and you decide to give someone a raise, you also pay for that yourself.”) In no-stakes conditions, they were told that they could simply decide that the ingratiation should get 5 extra guilders, and this would not affect their own compensation. (“After all, if you are the director of a company, and you decide to give someone a raise, you don’t pay for that from your own wallet.”)

In the observer conditions, the stakes were varied in the same way, but in this case, they pertained to whether the target of the description, not the subject, would have to give up money in order to assign more money to the ingratiation. So, no effects of stakes are to be expected in these conditions.

As in Experiment 1, the cover story was summarized and subjects were informed of what specifically the ingratiation had been told. The ingratiation’s impression description was the same as well, except for some changes related to the different cover story. For instance, the first sentence was, “I am Laura (Ronald) and I understand that they are getting 15 guilders and we are getting 5. That doesn’t seem fair to me, but anyway I’m not doing this for the money.” The last sentence was, “It’s such a shame that I will probably never meet her (him).”

Differing from Experiment 1, in Experiment 2 participants were asked, after reading the ingratiation’s description, to write down all of their thoughts in their own words on a blank sheet of paper; subsequently, they were asked how likeable, slimy, and sincere the ingratiation was, embedded within other trait scales. Participants also indicated on a 5-point scale (1 = *definitely not*, 5 = *definitely*) whether they felt like giving 5 guilders to the ingratiation, or, in observer conditions, whether they would if they were the one who had to make the decision.

Results and Discussion

Four participants, who were all in the high self-relevance condition (3 of them in the high-stakes condition), were discarded because they expressed suspicion in their responses to the last, open-ended question. Participant’s gender did not qualify the results—as in Experiment 1, judgments of females tended to be more favorable than of males, main effect, $F(3, 110) = 2.63, p < .06$, but interactions with gender were nonsignificant—and was dropped from the analyses below.

Participants were more willing to give extra money to the ingratiation when they did not have to give up their own money (3.76 vs. 2.68); main effect of stakes, $F(1, 116) = 23.74, p < .01$. This effect occurred regardless of whether participants were themselves the target or whether someone else was (interaction $F < 1$).

A 2 (stakes) $\times$ 2 (self-relevance) MANOVA on the likeability ratings, slime ratings, and ratings of whether the description was sincere produced a main effect of self-relevance, $F(3, 114) = 4.29, p < .01, \eta^2 = .10$. Compared with observers, participants who were the target of ingratiation judged the ingratiation as more likeable (5.35 vs. 4.35), $F(1, 116) = 12.90, p < .01$; less slimy (4.96 vs. 5.43), $F(1, 116) = 2.48, ns$; and more sincere (4.21 vs. 3.57), $F(1, 116) = 4.92, p < .05$. The main effect of stakes ($F < 1$) and the interaction, $F(3, 114) = 1.90$, were nonsignificant.

Participants’ open-ended thought listings were quantified by two independent judges. Disagreements were resolved by means of

discussion. Thoughts about the ingratiation were categorized as either favorable, unfavorable, or reflecting suspicion about the ingratiation’s motives. Analyses on these data did not reveal anything in addition to the results obtained from the closed-ended rating scales (e.g., there were more positive and less negative thoughts about the ingratiation in self-relevance than in control conditions). Also, judges coded whether participants noted that reading the description made them feel good; not surprising, this occurred far more frequently in self-relevance than in control conditions (.45 vs. .05), $F(1, 116) = 21.85, p < .01$. No other effects were found on this variable.

In addition, thoughts about the *target* of the description were categorized as indicating that the description was accurate (e.g., “Remarkable that his impression of me is so accurate” in self-relevance conditions; “Frances is apparently a really wonderful person” in control conditions) or inaccurate (e.g., “Some of the things she said about me don’t make sense” in self-relevance conditions; “I can’t believe that Frances is such a totally perfect person, everybody has negative things as well” in control conditions). Frequencies of these two categories ($r = -.04$) were submitted to a 2 (stakes) $\times$ 2 (self-relevance) $\times$ 2 (self-esteem, dichotomized at the median of 5.5) MANOVA. As noted previously, self-esteem had no effect on ratings of the ingratiation or other dependent variables; it did however produce an effect in this particular analysis. Specifically, a Self-Relevance $\times$ Self-Esteem interaction, multivariate $F(2, 111) = 3.87, p < .05$, on the number of times participants spontaneously noted that the description was *inaccurate*, univariate $F(1, 112) = 6.51, p < .05$, revealed that, in self-relevance conditions, participants with low self-esteem regarded the description as more inaccurate than those with high self-esteem (.28 vs. .06), $F(1, 53) = 7.91, p < .01$; in control conditions (i.e., when the description was not about participants themselves), this difference was absent (.03 vs. .10; $F < 1$). Thus, even though self-esteem did not qualify the effects of self-relevance in ratings of the ingratiation whatsoever, it did affect participants’ thoughts about the accuracy of the description. This issue will be readdressed in the General Discussion.

Experiment 3

Altogether, the target–observer difference turned out to be rather robust in these studies: The effect seemed relatively insensitive to variations in participants’ motives. Experiment 3 was conducted for three purposes. First, the procedure of Experiment 2 was repeated, but the (financial) stakes were increased, to examine if this would have more impact.

Second, in order to examine the possible role of mood, as discussed in the introduction, a mood scale was administered after the ingratiation episode. Participants’ thought listings in Experiment 2 already suggested that the flattering description of themselves made them feel good. As noted in the introduction, it is possible that the target–observer effect is mediated by mood. This would imply that targets have a better mood than observers (Hypothesis 1), and the target–observer effect is reduced when the effects of mood are accounted for (Hypothesis 2).

Third, in this experiment, an assessment was included of the time participants took to read the description by the ingratiation. This way, it could be examined if the condition with higher stakes for the participant produced longer processing, as was expected

(Hypothesis 3). In addition, if mood mediates the target–observer effect, the data on processing time are useful to examine if a more positive mood is associated with quicker—presumably more heuristic—processing. That is, targets may process more quickly and less elaborately than observers because the description puts them in a good mood and they want to avoid considering the possible ulterior motives of the ingratiation to maintain their positive mood (Hypothesis 4A). However, an opposite prediction about the effects of self-relevance on processing time is conceivable as well: Regardless of mood, targets may process more slowly and extensively than observers, simply because they are considering a description that is about themselves, and this self-relevant information is far more interesting and arouses more cognitive activity than a description about someone else (Hypothesis 4B).

Method

Participants were 115 freshman psychology students (83 women, 32 men), recruited during classes. They were paid 20 guilders (at the time, about \$10) for participating in this study and two independent studies used as fillers. The design and procedure were the same as in Experiment 2. In the personality test with which the experiment began, the scale for interpersonal orientation was dropped, and a scale for dominance orientation⁴ (from Goodwin, Gubin, Fiske, & Yzerbyt, 2000) was added.

Participants were told that whereas they would receive 20 guilders for their participation, the participants in the other group were taking part voluntarily and were not receiving any reward at all. They were told that they (or the target, in observer condition) would have the opportunity to give 10 guilders (about \$5) to the ingratiation, either using their own money or not, depending on the stakes condition (as in Experiment 2). In other respects, the story was the same as in Experiment 2, and so was the description given by the ingratiation. Unobtrusively, the time from appearance of this description on the screen up to the participant's pressing the *Return* key to move on was registered.

After the thought-listing task, which was the same as in Experiment 2, participants' mood was assessed by means of five positive and five negative adjectives: tense, excited, content, worried, annoyed, proud, delighted, confused, sad, happy. Participants were asked to indicate for each adjective, using a 7-point response scale (1 = *not at all*, 7 = *highly*), "how you feel at this particular moment as you're sitting here." Subsequently, dependent measures were assessed, participants were asked to indicate the likelihood that they (or the target, in observer conditions) would give 10 guilders to the ingratiation, and manipulation checks were administered. At the end of the experiment, participants in target conditions were asked directly whether they would assign 10 guilders to the ingratiation (yes or no), and were told that they would be committed to their response.

Results and Discussion

Five participants, 4 of whom were in the high self-relevance condition (and 2 in the high-stakes condition) who suspected that the ingratiation did not really exist, or was not really taking part without getting any reward, were excluded. As in the previous experiments, a main effect of gender, $F(3, 100) = 3.00, p < .05$, reflected more favorable judgments of women than men, but gender did not qualify the other effects and was discarded from subsequent analyses.

Giving Money to the Ingratiation

As in Experiment 2, participants were more willing to give money to the ingratiation when they did not have to give up their

own money than when they did (4.09 vs. 3.20); main effect of stakes, $F(1, 106) = 14.20, p < .01$, regardless of whether they were themselves the target or whether someone else was (interaction $F < 1$). A similar result emerged for the yes–no question at the end of the experiment, when it actually came down to giving the money or not; this question was put only to participants in the target condition, who were in fact giving away their own money in the high-stakes condition. In the no-stakes condition, only 2 out of 27 participants refused to give 10 guilders to the ingratiation; in the high-stakes condition, 13 out of 26 participants refused, $\chi^2(1, N = 53) = 12.85, p < .01$.

Judgments of the Ingratiation

A 2 (stakes) $\times$ 2 (self-relevance) MANOVA on ratings of likeability, sliminess, and sincerity, produced a main effect of self-relevance, $F(3, 104) = 4.51, p < .01, \eta^2 = .11$. Again, compared with observers, target participants judged the ingratiation as more likeable (5.60 vs. 4.54), $F(1, 106) = 12.57, p < .01$; less slimy (4.51 vs. 5.32), $F(1, 106) = 5.94, p < .05$; and more sincere (4.70 vs. 3.88), $F(1, 106) = 7.85, p < .01$. Regarding the effect of stakes, the interaction with self-relevance was nonsignificant for all three ratings (all $F < 1$), and only a main effect on slime ratings emerged, $F(1, 106) = 4.05, p < .05$. The ingratiation was judged as more slimy when there were no stakes involved for the target than when there were (5.25 vs. 4.60), regardless of whether the participant was the target or an observer. Possibly, this effect emerged because the flattery of the ingratiation was seen as relatively disproportionate in conditions where the target did not have to give up anything to reward the ingratiation. No other significant effects of stakes emerged.

Mood

Factor analysis on the mood items revealed one major factor explaining 37.3 % of the variance (eigenvalue 3.73), with positive loadings for the positive adjectives and negative loadings for the negative ones. After recoding the latter, Cronbach's α for the 10 items was .78. An ANOVA on the mean of this scale produced only a main effect of self-relevance, $F(1, 106) = 4.30, p < .05$. As expected (Hypothesis 1), targets of ingratiation were in a better mood than observers (5.22 vs. 4.89). However, although mood was significantly correlated with liking for the ingratiation ($r = .28$), slime ratings ($r = -.34$), and perceived sincerity ($r = .37$), it did not account for the self-relevance effect found on these variables: When mood was entered as a covariate in an ANOVA, the effect remained virtually the same, multivariate $F(3, 103) = 3.27, p < .05, \eta^2 = .09$; univariate F s for likeability, sliminess, and sincerity: $F(1, 105) = 9.45, p < .01$; 3.37, $p < .07$; and 4.74, $p < .05$, respectively.

⁴ Thanks to Stephanie Goodwin for making available the items of this scale on very short notice. Unfortunately, just like self-esteem, dominance orientation did not qualify the actor–observer effect, so it was not included in the analyses reported here.

Reading Time

The time it took participants to read the ingratiation's description⁵ was affected by self-relevance, $F(1, 106) = 4.87, p < .05$. Targets of flattery studied the description longer than observers (59.96 s vs. 50.62 s). There was a nonsignificant tendency for this effect to be qualified by stakes, $F(1, 106) = 3.61, p = .06$: The target–observer difference in reading time emerged only when the stakes were high (66.39 s vs. 48.78 s, $p < .05$), not when there were none (53.76 s vs. 52.53 s, *ns*). Duncan contrast tests showed that the mean of 66.39 in the self-relevance+stakes condition was significantly ($p < .05$) different from all three of the other means, which in turn did not differ from each other. That is, reading times increased only in the one condition in which participants were considering giving up their own money to the ingratiation (in all three other conditions, participants were not asked to give away their own money).

Thought Listings

Analyses on the open-ended thought listings produced similar results as in Experiment 2. Participants who were the target of ingratiation noted more frequently than control participants that the description made them feel good (.60 vs. .00), $F(1, 106) = 56.86, p < .01$; this effect was not qualified by any other variables. The correlation of this variable with mood was .36 ($p < .01$).

Also, as in Experiment 2, self-esteem only produced effects in the analyses on the number of times participants noted that the ingratiation's description was accurate or inaccurate ($r_{\text{accurate, inaccurate}} = -.12, ns$). A Self-Esteem (dichotomized at the median of 5.2) $\times$ Self-Relevance effect, multivariate $F(2, 101) = 3.15, p < .05$, was caused by the number of times noted that the description was accurate, univariate $F(1, 102) = 6.35, p < .05$. When participants were the target of the description, those with high self-esteem regarded the description as more accurate than those with low self-esteem (.38 vs. .08), $F(1, 49) = 5.97, p < .05$, whereas this difference was absent in control conditions (.11 vs. .06; $F < 1$). Differing from Experiment 2, the effect was nonsignificant ($F < 1$) for the number of times noted that the description was inaccurate.

In sum, the self-relevance effect was, again, replicated in this study; and, again, there was no interaction with stakes. Thus, as in Experiments 1 and 2, the target–observer effect is not affected by the target's motives—at least not by the motive to like the interaction partner (Experiment 1) or the motive to think twice because there are potential losses at stake (Experiment 2 and 3). Other results show that the stakes manipulation was effective in making participants think more elaborately: In target conditions, participants who could incur losses studied the ingratiation's description longer (confirming Hypothesis 3). The fact that reading times were longer only in this condition, in which participants' own resources were under consideration, indicates that the longer processing was not caused by the fact that participants were reading about themselves; if it were, reading times would have been longer in target+no-stakes conditions as well. Instead, we may assume that the longer processing was caused by the motive to form an accurate impression of the ingratiation, because tangible losses were at stake.

In spite of this apparent success of the stakes manipulation in making participants “think twice,” the target–observer effect was not affected by stakes. Remarkably, then, participants did not even reduce their liking for the ingratiation if only to rationalize their decision not to give him or her their money. (Possibly, they did not need to do this because there were other justifications, i.e., “All participants have already agreed to work for the money that they are getting” and “I’m not responsible for someone else who has agreed not to receive any money.”)

The present results also suggest that mood is not a crucial variable in the target–observer effect either. Although targets were in a better mood than observers, as predicted, this difference did not account for the differential judgments of the ingratiation.

Finally, as in Experiment 2, the open-ended thought listings showed that participants' self-esteem qualified the effects of self-relevance on only one type of variable, the number of times noted that the ingratiation's description of them was inaccurate (in Experiment 2) or accurate (in Experiment 3). This suggests that cognitive responses to flattery are affected by self-esteem, whereas affective responses are not. Indeed, the number of times participants mentioned that the description made them feel good was not affected by self-esteem at all.

Experiment 4

Experiments 1–3 found that the target–observer effect is not caused by possibly confounding variables in previous studies, namely, cognitive capacity, the motives that are aroused when forming an impression under (expected) interaction conditions, or mood. Although it is reassuring to establish that the target–observer effect is not an artifact, the absence of any qualifying variables in the previous studies also implies that we still do not know what really causes the effect. One possibility is that other motivational variables than the ones examined so far are at work. People may be motivated to accept flattery because it bolsters their self-esteem, even at the risk of losing tangible outcomes by trusting the ingratiation, or they may be motivated to reciprocate the liking by another person. In addition, it is possible that the target–observer difference is not primarily due to targets' favorable judgments of the ingratiation, but to observers' relatively less favorable judgments. Even though observers are the “control group,” it would be naive to regard them as having no stake in the matter and, hence, as not driven by any motivational biases. Observers probably bring their own biases into the equation. For instance, considering that people generally view themselves as better than average on social and likeability-related qualities (Allison, Messick, & Goethals, 1989), they may tend to dismiss extremely positive comments about others as excessive and, hence, ingratiation, assuming that another person could not possibly be that wonderful.⁶

In Experiments 1–3, observer participants read a description about someone they did not know. Assuming that the target was moderately likeable, like most people (cf. Skowronski & Carlston,

⁵ Distributions of reading times were inspected for outliers (i.e., three standard deviations from the mean), but there were none.

⁶ This possibility was suggested by an anonymous reviewer and by Constantine Sedikides.

1987), the ingratiation's positive comments may have seemed grossly overstated, which may have contributed to a less favorable judgment of the ingratiation. Participants in self-relevance conditions, on the other hand, read a positive description about themselves, which, presumably, was for the most part evaluatively consistent with what they were willing to believe about themselves. As a consequence, they may have had no reason to elaborate on the ingratiation's motives.

In Experiment 4, the role of the observer was examined more extensively. If observers are reluctant to accept extremely positive comments about someone who, for all they know, is not as likeable as they themselves are, this reluctance may be counteracted by giving them prior positive information about the target of ingratiation, containing qualities that are just as favorable as the qualities people tend to ascribe to themselves. In Experiment 4, this prior information contained all the characteristics that many students in a previous study (Vonk & Ashmore, 1993) ascribed to themselves in open-ended self-descriptions. Thus, the impression induced about the target was, though undeniably not as rich as the impression that was available for participants in self-relevance conditions, evaluatively similar. In sum, Experiment 4 has three conditions: target, observer, and observer+. This last label is used to refer to the new condition, where participants were observers but had additional, favorable knowledge about the target being described.

If, in this experiment, the observer and the observer+ conditions produce similar results, and both differ from the target condition, we may assume that the target–observer difference is caused primarily by motivational biases among targets of ingratiation and not by the absence of information about targets among observers, or by observers' motive to be cynical about excessive praise of someone else. If, on the other hand, the observer+ condition is similar to the target condition, this would suggest that observers' lack of favorable information about the target—information that people do have about themselves—is involved in the effect.

Method

Participants were 160 undergraduates (115 women, 45 men) with different majors. The procedure was the same as in Experiment 2, that is, participants were promised (and given) 15 guilders for participating and were told that the ingratiation was participating for 5 guilders; the target of the description (participants themselves in target conditions, someone else in observer and observer+ conditions) would be asked to give 5 guilders to the ingratiation, either from his or her own reward or not. Thus, a 3 (self-relevance: target, observer, observer+) $\times$ 2 (stakes) design was used.

There were four differences with Experiment 2. First, a minor difference was that in the first part of the study, the personality test battery—a test of narcissism—was added (Abridged Narcissistic Personality Inventory; Emmons, 1987).⁷ Second, as in Experiment 3, mood was assessed after participants read the ingratiation's description. Third, participants were not asked to list their thoughts after reading the description. Last but not least, the observer+ condition was added. This condition was identical to the observer condition, except that, shortly before reading the ingratiation's description of the target (Frances or Frank), participants were given a personality profile of the target. This profile was allegedly based on a previous study in which the target had played a "psychological game" with others; afterwards, everyone had to give a description of each other. In actuality, the description was based on an eyeball inspection of the open-ended self-descriptions from students collected in a previous study (Vonk & Ashmore, 1993); characteristics that were mentioned frequently in these

descriptions were selected for the profile of the target. Specifically, the description listed the following attributes: good sense of humor, cooperative, honest, opinionated, impatient, good common sense, someone with different sides—sometimes cheerful sometimes sad, sometimes serious, sometimes happy-go-lucky—has a lot of depth. Presumably, this description reflects to some extent the "modal self-concept" of undergraduates in the Netherlands. Participants were told that they would not really need this information but should regard it as background information. Subsequently, the procedure went on as in the other two conditions.

Results and Discussion

Six participants (5 of whom were in the high self-relevance condition, 3 in the high-stakes condition) were discarded because they suspected that the ingratiation did not exist or had been instructed to give this particular description of the target. In this experiment, gender did not have any significant effects.

As in the previous studies, in conditions where no losses were at stake, participants were more likely to assign extra money to the ingratiation, or to indicate that the target would (3.80 vs. 2.69), $F(1, 148) = 40.90$, $p < .01$; again, targets in no-stakes conditions more frequently assigned the money to the ingratiation (20 out of 24) than targets in high-stakes conditions (8 out of 23), $\chi^2(1, N = 47) = 12.07$, $p < .01$.

A 3 (self-relevance) $\times$ 2 (stakes) MANOVA on liking, slime ratings, and perceived sincerity did not produce any effects of stakes (but see Footnote 7 for effects of Stakes $\times$ Narcissism). The effect of self-relevance was significant, $F(6, 290) = 2.61$, $p < .05$, $\eta^2 = .05$. Table 1 shows the pertinent means, contrast tests, and univariate F statistics. As can be seen, the results for the observer+ condition are highly similar to those for the target condition, whereas both differ from the observer condition.

Reading times did not differ as a function of self-relevance or stakes (all $F < 1$). Because the target–observer difference in reading times emerged only in the high-stakes condition of Experiment 3, this suggests that the lower stakes in the present study (potential loss of 5 guilders instead of 10) were too low to bring about the previously obtained effect on attention.

The overall effect of self-relevance on mood was nonsignificant, $F(2, 148) = 2.41$, $p < .10$, but mood in the target condition was significantly ($p < .05$, two-tailed) more positive ($M = 5.24$) than in the observer and observer+ conditions (both 4.84). Thus, although the judgments of the ingratiation in the observer+ condition were very similar to the target condition, participants' mood in this condition was more similar to the observer condition. This result further corroborates the independence of self-relevance effects on mood versus judgment: Although participants' mood is positively affected when they are the target of ingratiation, the more favorable judgments of the ingratiation do not depend on this. Instead, the present results suggest that the effect emerges simply when

⁷ As the scales for self-esteem and dominance orientation, narcissism ($\alpha = .72$) did not produce any main effects or interactions with self-relevance, so it was dropped from the analyses reported here. It did, however, interact with stakes in judgments of likeability, $F(1, 142) = 4.29$, $p < .05$, and sliminess, $F(1, 142) = 5.23$, $p < .05$. Participants high on narcissism (after a median split) judged the ingratiation more negatively on all traits in high-stakes conditions (i.e., regardless if they were themselves the ones that could incur a loss or someone else was), whereas the converse applied to participants low on narcissism.

Table 1
Mean Judgments of Likability, Sliminess, and Sincerity of the Description

Condition	Likable ^a	Slimy ^b	Sincere ^b
Observer	4.84 _a	5.39 _a	3.75 _a
Observer +	5.50 _b	4.79 _b	4.40 _b
Target	5.96 _b	4.70 _b	4.43 _b
<i>F</i> (2, 148)	6.74	2.88	3.74
<i>p</i> <	.01	.06	.05

Note. Within each column, means with noncommon subscripts are significantly different (Duncan multiple range test, $p < .05$, two-tailed).

^a Scale anchors: 1 = *dislikable*, 7 = *likable*. ^b Scale anchors: 1 = *not at all*, 7 = *highly*.

there is an evaluative consistency of the ingratiation's comments with what participants believe or like to believe about the target (either the self or someone else).

General Discussion

In the present experiments, the target–observer effect was found to be remarkably robust and insensitive to other manipulations. In Experiment 1, the effect was not qualified by expected interaction; in Experiment 2, accuracy motivation (operationalized in terms of potential losses) did not reduce the effect; and in Experiment 3, no indication was found that the effect was mediated by mood. With regard to motivation, it should also be noted that other studies conducted in our laboratory, with different manipulations of accuracy motivation, have produced similar results (Jostmann, 2000). In these studies, participants were told that they would be financially rewarded if their impression of the ingratiation was accurate (Experiment 1), or they were told that the accuracy of their impression was informative about their social intelligence (Experiment 2). Neither of these manipulations affected the target–observer effect. In addition, it is interesting to note that none of the individual differences variables assessed in these studies appeared to be related to the effect: not self-esteem (assessed in all four experiments), not dominance orientation (assessed in Experiments 3 and 4), and not narcissism (assessed in Experiment 4). Hence, no evidence whatsoever was obtained that the target–observer effect is qualified by motivational or personality variables.

This conclusion implies that several possible causes of the effects, some of which follow from confounds in previous experiments, can be ruled out. First, the present results suggest that the target–observer effect is really due to differences in who the target of ingratiation is, and not to differences in the motive to like or to be accurate.

Second, the effect does not appear to be caused by vanity, as Jones (1964) and others (e.g., Kipnis & Vanderveer, 1971) have assumed, at least not if we conceive of vanity as a state of high self-esteem rather than a motive toward being liked and admired: Across four experiments, no evidence was found that the effect is stronger among participants with high self-esteem.

The results of Experiment 4 rule out another possible interpretation of the target–observer difference obtained in these studies, namely, that target participants accepted the flattery because it converged with their responses to the personality test on which the ingratiation's description of them was based. For instance, suppose

most participants responded to the personality items in a socially desirable way, trying to convey a favorable impression of themselves. In this case, they would expect a favorable description from the ingratiation, consistent with their responses to the test items. Hence, it would be rational to infer that the ingratiation was sincere. However, Experiment 4 produced the same results for targets as for observers who had a favorable impression of the target. Because these observers did not know how the target had responded to the personality test, they could only compare the ingratiation's description with their own impression, and not with the target's responses to the test items.

Fourth, the mood data obtained in Experiments 3 and 4 demonstrate that, although being flattered does put the target in a good mood, this effect is independent of the increased liking for the ingratiation: Mood did not mediate the target–observer difference. Also, the more positive mood did not result in less elaborate processing. Moreover, Experiment 4 showed that enhanced liking for the ingratiation can emerge without the more positive mood found in target conditions.

So what does cause the target–observer difference? Considering the results of Experiment 4, one possible explanation is that the effect is mainly cognitive. According to this interpretation, targets of flattery compare the ingratiation's comments with their own view of themselves; if the comments are consistent with their self-concept, they make, in Kelley's (1973) terminology, an entity attribution (in which they themselves are the entity). They infer that the ingratiation's statements are based on characteristics of the entity. As a consequence, they discount other possible causes of the behavior, such as the ingratiation's motives. Essentially, this process does not differ from other more or less rational entity attributions; for instance, when we observe that "John laughs at the comedian," and we already know that this comedian is very funny, we will attribute the behavior to the comedian (i.e., the entity) without considering dispositional causes on John's part. The only difference is that, in the present study, target participants are themselves the entity of the attribution. Because most psychologically healthy people, as the participants in these studies, have a positive view of themselves, they are likely to regard the flattery as evaluatively consistent with their self-concept, so that the ingratiation's motives are not questioned. In observer conditions, on the other hand, participants cannot compare the ingratiation's statements with their knowledge about the target, because they simply lack information about the target. As a consequence, in interpreting the behavior, they are more likely to consider the actor's motives, because it is all they have: They know the context in which the actor made the statements, but they do not know anything else. In this case, it is more likely that the actor's motives are questioned. In sum, this explanation assumes that, if the ingratiation's statements converge with what participants know or expect of the target, they infer that the description is favorable simply because it is the truth, and discount other variables. This happens regardless of whether the target is the self or someone else.

This view, however, is counterfeited by the results regarding self-esteem. If the fit between the ingratiation's statements and the participant's self-concept is crucial, then the target–observer difference should be reduced among participants with low self-esteem. In the present experiments, no evidence was found for this—and certainly not for lack of trying to find it. In addition to median split analyses with self-esteem, it was also examined whether there are differences between participants with high and

low self-esteem in a more absolute sense, that is, those scoring above versus below the midpoint of the 7-point self-esteem scale. This analysis was conducted on the combined data of all four experiments, so that there would be a sufficient number of participants with low self-esteem ($N = 55$ with self-esteem below 4, $N = 407$ with self-esteem of 4 and higher). Analyses with this variable did not produce any self-esteem effects either. Similarly, an analysis with three self-esteem categories (low, middle, high) failed to produce any effects. It was also examined whether the participants who were excluded in the studies because of their suspicion might be lower in self-esteem: As noted earlier, most of these participants were in the high self-relevance conditions, and they may have become suspicious precisely because the description did not match their self-concept. However, this was not the case either.

The one effect of self-esteem that did emerge in Experiments 2 and 3 was the interaction with self-relevance on the number of times participants noted that the description was accurate or inaccurate. High self-esteem participants who were the target of ingratiation noted more frequently than low-esteem participants that the description was accurate (Experiment 3), and low self-esteem participants more often felt that the description was inaccurate (Experiment 2). Note that these variables reflect primarily a cognitive process of comparing the ingratiation's description with what is known about the target. If the target-observer effect is cognitive in nature, then this process, and the conclusion that the ingratiation's description does or does not match one's self-concept, should affect subsequent judgments of the ingratiation. This, however, was not what happened. Whereas low self-esteem participants did perceive the description as less accurate, this did not affect their positive feelings and their judgments of the ingratiation at all. (Indeed, correlations between judgments of the ingratiation and the frequency with which participants noted that the description was accurate or inaccurate were all lower than .15 and nonsignificant.)

This pattern of findings converges with Swann, Griffin, Predmore, and Gaines's (1987) view, that people's cognitive responses to self-relevant feedback are motivated by the need for self-consistency, whereas their affective responses are motivated by self-enhancement. This independence between cognitive and affective responses is nicely illustrated by some of the thought listings from participants with low self-esteem. For instance, one participant in Experiment 2 wrote, "Huh? How can someone be so positive about me? Cool! I think she's a very very nice person." A participant in Experiment 3 wrote, "My first thought was: so many friendly things about me in one story, that can't be true. . . . I felt flattered. I thought it was very sweet of her to write this."

In sum, it appears that the cognitive responses to the flattering description (i.e., thoughts as to whether the description was accurate or not) did not affect judgments of the ingratiation. This conclusion was corroborated by analyses of covariance in which references to the description being accurate or inaccurate were entered as covariates: Neither among participants with high self-esteem nor among those with low self-esteem did these covariates affect the self-relevance effect on judgments of the ingratiation. Thus, we may conclude that these judgments are affected primarily by self-enhancement motives. Indeed, the present results support the notion that the motive to self-enhance is just as prevalent among low self-esteem as among high self-esteem participants (e.g., Sedikides & Strube, 1997; Swann et al., 1987). The fact that no single motivational or personality variable was found to qualify

the target-observer difference may then be interpreted as an indication of the strength of this motive. Of course, it is possible that other motives, or stronger manipulations of them, may reduce the effect, but for now we may conclude that the effect is quite robust and is independent of whether the flattery is seen as accurate or not.

As noted earlier, the results obtained for observers may be interpreted by the motive toward self-enhancement as well. Observers are not totally impartial and unbiased, and, in some way, their ego's are at stake too. Being motivated to assume, as most people are, that they are better than others, observers may be reluctant to uncritically accept lavish praise about another participant who just happened to be there at the same time. Ego-related considerations may have been especially salient in the present studies, where all participants were told initially that they would read a description about either themselves or someone else. Participants assigned to the control condition, who then got to read a description about someone else, may have wondered what the ingratiation would have said about them. They may have even felt some resentment that they were not the ones being flattered by the ingratiation. As a consequence, it would be very tempting to dismiss the flattery as a mere result of the ingratiation's dependence on the target, so that the description could be seen as uninformative. In Experiment 4, in the observer+ condition, this tendency may have been counteracted by the additional information about the target, showing that the target was in fact quite deserving of the praise she or he received.

This series of studies set out to identify the variables that could qualify or mediate the target-observer effect, yet it ended up scratching out each and everyone of them. Processing capacity, the motive to like one's interaction partner, the motive to avoid that one is being duped: No evidence was obtained that any of these variables qualifies the effect. This also goes for cognitive responses about the accuracy of the ingratiation's comments, which, as we saw earlier, were independent of judgments of the ingratiation, just as much as the affective variable mood was. This does not mean that these variables are not involved in the effect at all. Obviously, these are powerful variables that can affect judgments of others in many ways. Indeed, the present studies provided some illustrations of this: Expected interaction produced considerably more positive judgments; the possibility of incurring substantial losses led to increased processing; flattery resulted in a more positive mood among targets; and, independent of this, it made targets with low self-esteem question whether the ingratiation's comments were accurate. When people are flattered by others in everyday life, these variables are all confounded with each other, and they probably all contribute to targets' assessments of an ingratiation. The point is, however, that these variables do not have a mediating role over and above the simple fact that people like those who flatter them and tend to believe in their sincerity in spite of strong situational cues indicating otherwise.

References

- Allison, S. T., Messick, D. M., & Goethals, G. R. (1989). On being better but not smarter than others: The Muhammad Ali effect. *Social Cognition*, 7, 275-295.
- Backman, C. W., & Secord, P. F. (1959). The effect of perceived liking on interpersonal attraction. *Human Relations*, 12, 379-384.
- Baron, R. A. (1986). Self-presentation in job interviews: When there can be

- "too much of a good thing." *Journal of Applied Social Psychology*, 16, 16–28.
- Belmore, S. M. (1987). Determinants of attention during impression formation. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 13, 480–489.
- Berscheid, E., Graziano, W., Monson, T., & Dermer, M. (1976). Outcome dependency: Attention, attribution, and attraction. *Journal of Personality and Social Psychology*, 34, 978–989.
- Bless, H., Mackie, D. M., & Schwarz, N. (1992). Mood effects on attitude judgments: Independent effects of mood before and after message elaboration. *Journal of Personality and Social Psychology*, 63, 585–595.
- Bodenhausen, G. V. (1993). Emotions, arousal and stereotypic judgments: A heuristic model of affect and stereotyping. In D. M. Mackie & D. L. Hamilton (Eds.), *Affect, cognition, and stereotyping: Interactive processes in group perception* (pp. 13–37). San Diego, CA: Academic Press.
- Byrne, D., & Ramey, R. (1965). Magnitude of positive and negative reinforcements as a determinant of attraction. *Journal of Personality and Social Psychology*, 6, 884–889.
- Darley, J. M., & Berscheid, E. (1967). Increased liking as a result of the anticipation of personal contact. *Human Relations*, 20, 29–39.
- Ditto, P. H., & Lopez, D. F. (1992). Motivated skepticism: Use of differential decision criteria for preferred and nonpreferred conclusions. *Journal of Personality and Social Psychology*, 63, 568–584.
- Emmons, R. A. (1987). Narcissism: Theory and measurement. *Journal of Personality and Social Psychology*, 52, 11–17.
- Erber, R., & Fiske, S. T. (1984). Outcome dependency and attention to inconsistent information about others. *Journal of Personality and Social Psychology*, 47, 709–726.
- Fein, S. (1996). Effects of suspicion on attributional thinking and the correspondence bias. *Journal of Personality and Social Psychology*, 70, 1164–1184.
- Fein, S., & Hilton, J. L. (1994). Judging others in the shadow of suspicion. *Motivation and Emotion*, 18, 167–198.
- Fein, S., Hilton, J. L., & Miller, D. T. (1990). Suspicion of ulterior motivation and the correspondence bias. *Journal of Personality and Social Psychology*, 58, 753–764.
- Fodor, E. M., & Farrow, D. L. (1979). The power motive as an influence on use of power. *Journal of Personality and Social Psychology*, 37, 2091–2097.
- Forgas, J. P. (1998). Happy but mistaken? Mood effects on the fundamental attribution error. *Journal of Personality and Social Psychology*, 75, 318–331.
- Forgas, J. P., & Bower, G. H. (1987). Mood effects on person perception judgments. *Journal of Personality and Social Psychology*, 53, 53–60.
- Gilbert, D. T., Pelham, B. W., & Krull, D. S. (1988). On cognitive busyness: When person perceivers meet persons perceived. *Journal of Personality and Social Psychology*, 54, 733–740.
- Godfrey, D., Jones, E. E., & Lord, C. (1986). Self-promotion is not ingratiating. *Journal of Personality and Social Psychology*, 50, 106–115.
- Goodwin, S. A., Gubin, A., Fiske, S. T., & Yzerbyt, V. (2000). Power can bias impression processes: Stereotyping subordinates by default and by design. *Group Processes and Intergroup Relations*, 3, 227–256.
- Gordon, R. A. (1996). Impact of ingratiation on judgments and evaluations: A meta-analytic investigation. *Journal of Personality and Social Psychology*, 71, 54–70.
- Jones, E. E. (1964). *Ingratiation*. New York: Appleton-Century-Crofts.
- Jones, E. E. (1990). *Interpersonal perception*. New York: Freeman.
- Jones, E. E., Gergen, K. J., & Jones, R. G. (1963). Tactics of ingratiation among leaders and subordinates in a status hierarchy. *Psychological Monographs*, 77(Whole No. 566), 1–20.
- Jones, E. E., Stires, L. K., Shaver, G. K., & Harris, V. A. (1968). Evaluation of an ingratiation by target persons and bystanders. *Journal of Personality*, 36, 349–385.
- Jostmann, N. (2000). *The target-bystander effect*. Unpublished master's thesis, University of Nijmegen, Nijmegen, the Netherlands.
- Kelley, H. H. (1973). The process of causal attribution. *American Psychologist*, 28, 107–128.
- Kipnis, D., & Vanderveer, R. (1971). Ingratiation and the use of power. *Journal of Personality and Social Psychology*, 17, 280–286.
- Kunda, Z. (1990). The case for motivated reasoning. *Psychological Bulletin*, 108, 480–498.
- Neuberg, S. L., & Fiske, S. T. (1987). Motivational influences on impression formation: Outcome dependency, accuracy-driven attention, and individuating processes. *Journal of Personality and Social Psychology*, 53, 431–444.
- Pyszczynski, T. A., & Greenberg, J. (1987). Toward an integration of cognitive and motivational perspectives on social inference: A biased hypothesis-testing model. *Advances in Experimental Social Psychology*, 20, 297–341.
- Rosenberg, M. (1965). *Society and the adolescent self-image*. Princeton, NJ: Princeton University Press.
- Sedikides, C., & Strube, M. (1997). Self-evaluation: To thine own self be good, to thine own self be sure, to thine own self be true, and to thine own self be better. *Advances in Experimental Social Psychology*, 29, 209–269.
- Skowronski, J. J., & Carlston, D. E. (1987). Social judgment and social memory: The role of cue diagnosticity in negativity, positivity, and extremity biases. *Journal of Personality and Social Psychology*, 52, 689–699.
- Swann, W. B., Griffin, J., Predmore, S., & Gaines, B. (1987). The cognitive-affective crossfire: When self-consistency confronts self-enhancement. *Journal of Personality and Social Psychology*, 52, 881–889.
- Tjosvold, D. (1978). Affirmation of the high-power person and his position: Ingratiation in conflict. *Journal of Applied Social Psychology*, 8, 230–243.
- Tyler, T. R., & Sears, D. O. (1977). Coming to like obnoxious people when we must live with them. *Journal of Personality and Social Psychology*, 35, 200–211.
- Vonk, R. (1994). Trait inferences, impression formation, and person memory: Strategies in processing inconsistent information about persons. In W. Stroebe & M. Hewstone (Eds.), *European review of social psychology* (Vol. 5, pp. 111–149). New York: Wiley.
- Vonk, R. (1998a). Effects of cooperative and competitive outcome dependency on attention and impressions. *Journal of Experimental Social Psychology*, 34, 265–288.
- Vonk, R. (1998b). The slime effect: Suspicion and dislike of likeable behavior towards superiors. *Journal of Personality and Social Psychology*, 74, 849–864.
- Vonk, R. (1999a). Differential evaluations of likeable and dislikeable behaviours enacted towards superiors and subordinates. *European Journal of Social Psychology*, 29, 139–146.
- Vonk, R. (1999b). Effects of outcome dependency on correspondence bias. *Personality and Social Psychology Bulletin*, 25, 110–117.
- Vonk, R. (2001). Aversive self-presentations. In R. M. Kowalski (Ed.), *Behaving badly: Aversive interpersonal behaviors* (pp. 79–155). Washington, DC: American Psychological Association.
- Vonk, R., & Ashmore, R. D. (1993). The multi-faceted self: Androgyny reassessed by open-ended self descriptions. *Social Psychology Quarterly*, 56, 278–287.

Received November 1, 2000

Revision received September 10, 2001

Accepted September 25, 2001 ■