
CURRENT RESEARCH IN SOCIAL PSYCHOLOGY

<http://www.uiowa.edu/~grpproc/crisp/crisp.html>

Volume 10, Number 7

Submitted: November 16, 2004

First Revision: December 7, 2004

Accepted: December 8, 2004

Published: December 8, 2004

CONVERSATION LOGIC EFFECTS IN THE MINIMAL GROUP PARADIGM: EXISTENT BUT WEAK

Hartmut Blank

University of Leipzig, Germany

ABSTRACT

According to a conversation logic (Grice, 1975) analysis of the minimal group paradigm, participants discriminate along group boundaries because they feel obliged to incorporate the provided group membership information into their resource allocation decisions. Conversely, intergroup bias might disappear if the relevance of this information is explicitly attributed to a different task, as first demonstrated by Blank (1997). Two experiments addressing possible alternative interpretations of my earlier results, however, failed to support this expectation. In retrospect, the manipulation of group membership relevance might have been overridden by a counteracting salience effect. In contrast, a third experiment provided support for the conversation logic-based prediction that under some conditions outgroup bias should occur. Overall, however, conversation logic effects seem to be weak, compared to other influences in the minimal group paradigm. The general discussion focuses on the inherent uncertainty of the experimental setting and the heterogeneity of behavioral strategies it induces.

INTRODUCTION

Tajfel and associates developed the minimal group paradigm (MGP) to explore the minimal conditions for intergroup discrimination to occur (Tajfel, 1970; Tajfel, Billig, Bundy & Flament, 1971). In essence, the procedure consists of first categorizing participants into two arbitrary groups on the basis of a trivial criterion (thereby creating minimal groups in the sense that all features normally associated with group membership are missing: face-to-face interaction, common history, personal acquaintance, role structure, group norms, group cohesion, etc.) and then requiring individual participants to judge anonymous other ingroup and outgroup members on evaluative dimensions or to distribute rewards (money or points) between them. Typically, the participants' averaged evaluations or reward allocations exhibit ingroup bias, that is, on average, participants treat their own group more positively than the outgroup (see reviews by Brewer, 1979; Diehl, 1990; Mullen, Brown & Smith, 1992; Turner, 1978).

Tajfel and Turner (1986) have proposed Social Identity Theory (SIT) as an explanation of this phenomenon: People's group memberships contribute in important ways to their identities, and they also seek to derive positive self-esteem from the groups they are associated with. In the MGP, although the group categorization is trivial, its salience might induce participants to at least temporarily identify with it and try to make it as gratifying as possible, namely, by positively differentiating the ingroup from the outgroup. The salience aspect with respect to group identification has been emphasized particularly by self-categorization theory (Turner, 1985; Turner et al., 1987), a further development of SIT. Other developments within this theoretical tradition have highlighted the uncertainty reduction function of group categorization and subsequent ingroup bias (as opposed or in addition to self-enhancement; e.g., Abrams & Hogg, 1988).

However, there have also been critics of this theorizing and the underlying research paradigm (see, e.g., Rabbie, Schot, & Visser, 1989; and Schiffmann & Wicklund, 1992; for general critiques). Particularly important for the present work, some researchers have been concerned about possible demand characteristics of the MGP (Berkowitz, 1994; Gerard & Hoyt, 1974). The demand characteristics argument holds that participants do not show ingroup bias out of a self-enhancement motivation but are seduced to do so by perceived demand characteristics of the experimental situation. This refers primarily to the fact that the group membership of the participants to be rewarded in an allocation task is made exceedingly salient and may therefore lead the participants to think that the experimenter wants them to use this information and discriminate against the outgroup. Although a study by St.Claire and Turner (1982) seemed to have ruled out such an explanation, it has been revived by additional empirical evidence provided by Berkowitz (1994). Berkowitz found that participants clearly perceived demand characteristics towards ingroup bias and acted in accordance with these. Particularly interesting are postexperimental reports on the transmitters of the demand characteristics: Half of the participants indicated that the group membership information provided along with the distribution matrices had conveyed them the experimenter's likely hypothesis, and a third of the participants mentioned the categorization procedure. Thus it seems that, at present, demand characteristics cannot be excluded in the MGP.

However, some problems remain with the demand characteristics account as well. Firstly, although the salience of the group membership information somehow seems to convey to the participants that the experiment has something to do with intergroup behavior, it is not immediately clear why this would lead them to suspect that ingroup-favoring behavior (and not outgroup-favoring behavior or fairness) is expected of them. To some degree, there is an inherent circularity in the demand characteristics argument, because it presupposes to some degree the phenomenon it tries to explain. That is, the fact that participants seem to perceive directional demands (as in Berkowitz, 1994) may in itself be a consequence of their experiences with intergroup phenomena (e.g., knowledge about social norms or social identity processes). At least, it seems that some additional factor is needed to account for the direction of the expected and displayed discrimination (e.g., social norms, see below).

Secondly, while specific situational cues (the categorization procedure and in particular the group membership information in the distribution matrices; see above, Berkowitz, 1994) seem to be important in shaping participants' expectations about the experiment, the exact mechanism by which this is brought about is not specified. The latter, in turn, would allow critics to even question their causal role within the demand characteristics process. For example, one might argue that not the cues transfer the demand characteristics upon which the participants act in the experiment, but instead the participants, when asked in the post-experimental questionnaire, point to these salient antecedent features as convenient rationalizations of their own behavior.

In a sense, the demand characteristics account is unsatisfactory because it provides only a description of the participants' perceptions in (more precisely, after) the experiment (which are then offered as an explanation of the participants' behavior) and does not explain how these perceptions originate in the experimental situation.

As an approach that might serve to bridge this theoretical gap, I (Blank, 1997) have proposed a partial account of the ingroup bias phenomenon in the MGP based on Grice's (1975) conversation logic. This approach provides a mechanism that would explain why the participants draw on the group membership information in order to guide their behavior in the experiment. However, it makes no predictions regarding the precise manifestation of intergroup behavior (i.e., ingroup or outgroup favoritism). Insofar, it is only a partial account for the usual MGP findings. Nevertheless, it yields predictions that are tested in the three experiments reported below.

The general rationale is that the same pragmatic rules of communication operate in psychological experiments as in everyday conversations and this must be taken into account when interpreting the results of such experiments lest one runs the risk of serious misinterpretations. Particularly, each step in an experimental procedure (e.g., instructions, presentation of information, measurements, repetition of measurement, etc.) constitutes a communicative act which is interpreted by the participant on the basis of certain conversational rules or, in Grice's terminology, maxims. Importantly, when analyzed before the communicative background of the experiment, experimental manipulations may take on different meanings from what was intended by the experimenter, opening the door for alternative interpretations of the investigated phenomena in terms of conversation logic. Such accounts have recently been provided in such different fields as attribution, eyewitness testimony, judgement and decision, and opinion surveys (Blank, 1998; Bless, Strack & Schwarz, 1993; Fiedler, 1988; Hertwig & Gigerenzer, 1999; Hilton, 1995; Schwarz, 1999).

Applied to the MGP, the logic of conversation approach concentrates particularly on the group membership information given in the reward allocation task. However, unlike the demand characteristics account, it provides a mechanism by which the participants come to use this information. According to Grice's maxim of relevance, each contribution to a conversation is perceived to be relevant to the topic of the conversation. Thus, when information about the group membership of two persons to be rewarded is given along with the distribution matrices, participants will assume that this information is relevant for the task at hand (otherwise it would not be presented) and therefore feel that they should somehow take it into account when making their allocation decisions (moreover, it is the only useful information the participants have to guide their allocation decisions, which should enhance reliance upon it). Insofar as this information is one about a difference (in the most interesting case of one ingroup and one outgroup member to be rewarded), the default reaction mode to this information should be to make a difference, namely, in the rewards allocated to the ingroup and outgroup members.

Of course, this is not to say that each and every participant has to and will use the group membership information in this way. Some people may recognize that the group membership information may be relevant for the task (or, that the experimenter wants them to perceive it as relevant), may also clearly realize the option of making a difference, and may nevertheless decide to ignore this option and allocate the rewards according to other principles (e.g., equality). Indeed, human behavior is expected to display some variability. Still, this does not invalidate the present approach, because firstly there would be observable effects at the group level even if only some participants would "fall prey" to the conversational mechanism advocated here, and secondly people's behavior in the MGP itself is known to be extremely variable. For example, there is typically a sizeable proportion of participants who distribute fairly (Turner, 1983).

Of more importance is an inherent limitation of the present approach (which, however, allows for additional predictions; see Experiment 3): "Making a difference" does not predict the direction of the difference (i.e., pro ingroup or pro outgroup), and therefore the conversation logic account can only be a partial explanation of ingroup bias. For a complete account of ingroup bias, this explanation must be supplemented by other, ingroup-favoring mechanisms (e.g., self-enhancement via positive distinctiveness as posited by SIT, or a generic norm of loyalty to the ingroup as initially suggested by Tajfel et al., 1971; see also Gaertner & Insko, 2001; Hertel & Kerr, 2001; for the impact of social norms in the MGP). Nevertheless, because it disentangles differentiation and direction, the conversation logic approach suggests a unique experimental approach (beyond the assessment of perceived demand characteristics), as will be detailed below.

To reiterate, the conversation logic approach proposed here holds that the presentation of information about an ingroup-outgroup difference and the perception of this information as relevant for the allocation task are necessary (but not sufficient) conditions for ingroup bias to occur. With an eye towards experimentation, this specification of a relevance perception process as a precondition for ingroup bias allows us to identify conditions under which ingroup bias might disappear. If the group membership information is presented but not perceived as relevant for the task, then the participants will not feel obliged (by virtue of adherence to cooperative communication principles) to use it, and therefore no ingroup bias might result. Such a state of affairs could be achieved if the participants had an alternative possibility to attribute the relevance of the group membership information, for instance, a second task besides the reward allocation task for which this information is explicitly made relevant.

This was the idea I explored in a previous study (Blank, 1997), which otherwise followed the usual MGP. That is, the participants were categorized into two minimal groups and subsequently distributed rewards between ingroup and outgroup members. Half of the participants, however, worked on a second task in combination with the reward allocation task. They were requested to remember, after three matrices each, the points they had given to each member of the ingroup or outgroup, a task for which the group membership information was obviously relevant (however, the relevance of this information to the allocation task was not explicitly denied). In order to prevent easy shortcuts in solving this task, several numerically different versions of the distribution matrices were constructed. The basic result was that ingroup bias was absent in the double-task group, whereas the usual ingroup bias could be replicated in the standard group, which is in accordance with the conversation logic approach.

Yet these results were equivocal because there are at least two alternative interpretations (brought up by a reviewer of the 1997 article): First, the secondary memory task was relatively difficult and demanding, so that the participants might have concentrated on this task at the expense of the allocation task (plainly speaking, they were too busy to discriminate). Second, the memory task had unequivocal and easily checkable solutions and therefore probably induced evaluation apprehension in the participants. Because it was also a difficult task, mastery of it might give them an opportunity to present themselves favorably to the experimenter and thereby enhance self-esteem in way that bypasses possible self-enhancement through ingroup bias (i.e., they might show interpersonal instead of intergroup behaviour; Tajfel & Turner, 1986).

The first two experiments presented here sought to test the same conversation logic prediction as above while systematically exploring the impact of these alternative mechanisms. In Experiment 1, the difficulty of the secondary task as well as the relevance of the group membership information for this task were experimentally manipulated, with the conversation logic expectation that ingroup bias should disappear at either level of task difficulty, provided that the group membership information is relevant for it. Both alternative interpretations would predict that a difficult task suppresses ingroup bias regardless of the perceived relevance of the group membership information for it. In Experiment 2, the relevance issue was investigated even more directly and without a secondary task by plainly telling half of the participants that the group membership information was irrelevant for the reward allocation task but would be needed in a later task. According to the conversation logic approach, no ingroup bias should result. Finally, Experiment 3 tested a new prediction of the conversation logic approach. As argued above, the perceived relevance of the group membership information only prompts the participants to make a difference but does not per se imply the direction of this difference. Consequently, it should also be possible to systematically induce outgroup favoritism under suitable conditions.

EXPERIMENT 1

Method

Participants

Hundred and forty-two students from various disciplines except psychology took part in the study. They were recruited from diverse classes held at the campus to participate in "three psychological experiments" (see below) in exchange for a remuneration of 5 Euros. The data of 14 participants could not be analyzed because they had misunderstood the instructions or did not complete all matrices, etc. The remaining 128 participants had been randomly assigned to four experimental conditions (see below) with the restriction that (a) 32 of them participated in each condition, (b) counterbalancing was preserved within conditions (see below), and (c) within one experimental session only one condition could be realized. The number of participants within experimental sessions ranged from one to fifteen.

Procedure and Design

The first part of the study (designated "Experiment 1") served to categorize the participants into two minimal groups with the help of an ostensible colour perception test which required them to make five choices, on five-point scales, between pairs of colours placed successively on a sheet of paper, according to their preference for one or the other colour. The experimenter [1] collected the finished "test sheets", and while the participants worked on a 15-minute filler task (designated "Experiment 2"; a study on autobiographical memory), he pretended to calculate each participant's individual colour perception "test result" with the help of a computer notebook. In fact, the experimenter did not calculate individual scores but randomly informed half of the participants that they were of the "colour sensitive type" or the "contrast sensitive type," respectively. This feedback was embedded in the written instructions to the reward distribution task ("Experiment 3"). The instructions announced the reward distribution task as a decision making task that required to have two groups of participants. For convenience, the two types of perceivers as identified in the first "experiment" (which, so the participants learned, were about evenly distributed in the population) would be used for this purpose. Then each participant read that he or she was a member of the "colour sensitive" or "contrast sensitive" group. Their task would be to distribute reward points between two anonymous people identified only by their participation number and group membership. A filled-in example matrix (showing mild ingroup bias) followed. Further, we emphasized that it was not possible to allocate rewards to oneself. However, as an incentive, we announced that the three participants with the highest sum of points awarded by the other participants each would win 10 Euro.

Depending on the experimental condition, additional instructions followed with respect to the secondary task. Specifically, there were four experimental conditions: (1) The standard condition proceeded just as explained above, without a secondary task. (2) In the difficult only condition, the participants had to remember, after three matrices each, the participant numbers[2] of the persons to be rewarded (these numbers also figured in the matrices) as well as their position in the matrix (top row vs. bottom row; see matrices and dependent measures section). (3) The relevant-difficult condition was identical to the difficult only condition except that the participants had to remember the participant numbers and the group membership of the rewarded persons. (4) In the relevant-easy condition, the participants had a fairly easy secondary task: They were to remember, after each matrix, only the group membership of the persons in the matrix. In conditions 2 to 4, we explained the respective secondary tasks using an example, and instructed the participants to devote equal effort to both tasks.

Thereafter, the participants began with the reward allocation (plus secondary) task. When finished, they answered a final post-experimental questionnaire which asked for their distribution strategies, the perceived purpose of the experiment, the impact of the provided group membership of the persons, and reasons for ingroup bias or, alternatively, fairness in the distribution of rewards (I return to some relevant results from this questionnaire in the discussion section). Upon termination of the study, the participants were fully debriefed.

Matrices and Dependent Measures

Each matrix consisted of two rows. Each row represented possible payoffs for one person and first indicated on the left side the participant number (a one- or two-digit number) and group membership of the person. The possible payoffs for the persons in the two rows depended on the matrix type. I used three types of point distribution matrices to measure the prevalence of various distribution strategies. (1) A simple INFAV matrix (as employed in Tajfel et al., 1971, Exp. 1) assessed the degree to which the participants favored their own group vs. the outgroup in the reward allocations. (2) A MIP & MJP vs. MD matrix (as used by Tajfel et al. in their second experiment) measured the joint impact (or, "pull") of the two distribution strategies maximum ingroup payoff (MIP) and maximum joint payoff (MJP) on another strategy, maximum difference (in favor of the ingroup; MD). (3) The third matrix type was a variation of another commonly used matrix, F (fairness) vs. MIP & MD (e.g., Billig & Tajfel, 1973). This variation consisted of adding MJP to the fairness side of the matrix, thus contrasting the joint impacts of two non-discriminatory distribution strategies - F and MJP - and two discriminatory strategies - MIP and MD - on each other. For details on various distribution strategies and the logic of their assessment via pull scores, see Tajfel et al. (1971), Blank (1997), Bornstein et al. (1983), or Bourhis, Sachdev and Gagnon (1994).

In agreement with the proceeding in Blank (1997), I used eight versions of each matrix type in the present experiments. Four versions each differed in the numerical values of the rewards to be distributed, although they obeyed the same construction principle. For example, in a standard version of an ingroup favoritism matrix, the values in the top row (assigned to, say, an ingroup member) run from 1 to 14 whereas the bottom values (assigned to an outgroup member) run from 14 to 1. Then, a numerical variation of this principle would have, for instance, the top row running from 5 to 18 and the bottom row from 18 to 5. Another variation would have values from 2 to 28 in the top row and from 28 to 2 in the bottom row (however, no matrix contained any negative values in the experiments reported here). Also, these matrices differed from those conventionally used in that they consisted of only seven columns instead of thirteen or fourteen (making it easier to construct numerically different versions). Further, the four numerically different matrix versions were used in four different combinations of ingroup and outgroup members in the top and bottom rows of the matrices (i.e., ingroup top/ingroup bottom, ingroup top/outgroup bottom, outgroup top/ingroup bottom, and outgroup top/outgroup bottom). This served to counterbalance the assignment of numerical versions to member combinations across participants. Finally, each of the four versions of each matrix type had an additional mirror version with reversed right-left ordering of the points, yielding eight versions of each matrix type and 24 matrices altogether. Another three matrices placed at the beginning of the matrix booklets served as practice trials and were not analyzed.

Because the specifics of these various matrix types and distribution strategies are not a central focus of this article, I also calculated a more general measure as suggested by Diehl (1990), that is, the difference between the total points assigned to ingroup members and the total points assigned to outgroup members, summed across all 24 analyzed matrices. This measure will serve as the dependent measure throughout this article. The results based on this analysis are entirely consistent with the more specific analyses of the matrices as outlined above (these specific results may be obtained upon request from the author).

Results & Discussion

Table 1 gives the essential results of Experiment 1. In the standard condition, the usual ingroup bias effect could be replicated with one-tailed testing, the effect size (Cohen's *d*) being in the small to medium range, corresponding to the usual order of magnitude in minimal group experiments (Mullen et al., 1992). Contrary to my expectations, this ingroup bias effect did not vanish with a secondary task for which the group membership information was explicitly relevant. This held for both the relevant-difficult and relevant-easy conditions, where significant ingroup bias emerged. Another unexpected result was that ingroup bias was essentially absent in the difficult-only condition where the participants worked on a difficult secondary task for which the group membership information was not relevant.

Table 1: Average Ingroup Bias in Different Experimental Conditions of Experiment 1 (Difference between Total Points Assigned to the Ingroup and Outgroup)

Condition	Difference	SD	t(31)	p[a]	d
1. Standard	25.4	77.2	1.86	.04	.33
2. Difficult only	7.3	59.0	.70	.24	.12
3. Relevant-difficult	21.8	61.9	1.99	.03	.35
4. Relevant-easy	45.0	77.9	3.27	.001	.58

[a] One-tailed tests against the null hypothesis of no discrimination.

This latter result is perfectly in line, however, with both alternative interpretations of the Blank (1997) results as outlined in the introduction. That is, the difficult secondary task may have absorbed the cognitive resources of the participants, keeping them too busy to discriminate, and/or may have provided them with an opportunity to directly enhance their individual self-esteem, thereby obviating the need for indirect self-enhancement via ingroup favoritism.

Yet some aspects of the data cannot be fully explained by these alternative accounts either. Because, according to these accounts, the difficulty of the secondary task is responsible for a suppression of ingroup bias, they cannot explain why there is significant ingroup bias in the relevant-difficult condition. Also, they cannot explain why the amount of ingroup bias is roughly twice as high in the relevant-easy condition, compared to the standard condition, while there should be no difference - from these perspectives - between these conditions. Thus, the pattern of results creates problems for all of the previously discussed explanatory approaches.

A more convincing, post hoc interpretation of the pattern of results in Experiment 1 would result from the supposed operation of two separate principles: First, the difficulty of the secondary task serves to reduce ingroup bias, as expected from the two alternative interpretations of the Blank (1997) results in terms of cognitive load and direct self-enhancement. Second, the use of the group membership information in the two relevance conditions leads to increased salience of the group categorization, which in turn is known to enhance ingroup bias (e.g., Brewer, 1979). Such increased salience was also reflected in some participants' comments in the post-experimental questionnaire (e.g., "[I] only attended [to the group membership information] because it was later tested (...) If this had not been the case, I probably would have paid no attention to it"; statement by a participant in the relevant-easy condition).

Importantly, this effect of increased category salience would perfectly counteract the effect of the relevance manipulation that followed from my conversation logic analysis. Rather than freeing the participants from the (conversationally implied) demand to use the group membership information for discrimination, its stated relevance for the other task seems to have seduced at least some participants to use it as a guideline for their reward allocation decisions, this latter effect being stronger in hindsight. In retrospect, then, the four conditions realized in Experiment 1 constitute a 2 (secondary task load) * 2 (category salience) between-participants design in which the experimental conditions can be identified as follows: Standard = no load, low salience; difficult only = high load, low salience; relevant-difficult = high load, high salience; relevant-easy = (essentially) no load, high salience.

Having identified this post hoc design, it may be appropriate to conduct a post hoc ANOVA in order to assess the impact of secondary task load and category salience (treated as a random and a fixed factor, respectively) on the amount of ingroup bias. This ANOVA revealed marginally significant impacts of both factors (secondary task load: $F(1,1) = 63.23$, $p = .08$; category salience: $F(1,1) = 43.27$, $p = .10$). There was no significant interaction between these factors ($F < 1$).

Taken together, this analysis lends some support to the above post hoc interpretation of the Experiment 1 results. With respect to the original issue being investigated in this experiment, namely, the possible contribution of conversation logic mechanisms to ingroup bias in the minimal group paradigm, it seems then that the idea of manipulating the perceived relevance of the group membership information for the reward allocation task did not work very well because the participants' reward distributions were more thoroughly affected by two unintended side effects of this manipulation, namely, effects related to task difficulty and category salience. In particular, it seems that the obviously stronger but contrary effect of category salience on ingroup bias made it impossible to detect a relevance effect as expected from the conversation logic analysis.

There is, however, at least a single proof of existence for such a mechanism in Experiment 1, stemming from the postexperimental questionnaire. One participant (in the relevant-easy condition), when asked how the information about the group membership of the to-be-rewarded persons influenced his or her strategy in the distribution task, answered: "Not at all! K and F [the German initials of the categories] had a meaning only for the memory task". While this single statement certainly constitutes no impressive evidence for the conversation logic account, it points to the possibility that one might find more support for it with a different procedure that avoids the problems of the secondary task manipulation in Experiment 1. I tried this in Experiment 2.

EXPERIMENT 2

The key idea in Experiment 2 was to manipulate the perceived relevance of the group membership information for the reward allocation task without the help of a secondary task, in order to circumvent the problems associated with such a task, as detailed above. I did this by simply telling the participants in an irrelevance condition that the group membership information was not needed for the present reward allocation task but for another experiment that would be done later with the same materials. A standard condition identical to the one realized in Experiment 1 (except for minor changes due to the computer-controlled administration in Experiment 2) served as a control condition. The expectation from the conversation logic approach was that participants in the irrelevance condition should not feel obliged to use the group membership information to guide their reward allocation decisions, and therefore they should not exhibit ingroup bias in their allocations.

Method

Participants

Fifty-two psychology undergraduates participated in Experiment 2 in exchange for a payment of 5 Euro or 7.5 Euro, depending on whether they also participated in Experiment 3 (see below), or for equivalent course credit. All of them knew at the outset that some of them would be chosen randomly to participate in a second, shorter session (Experiment 3). By random assignment, 24 individuals in Experiment 2 participated in the standard condition and 28 in the irrelevance condition (with the restriction that counterbalancing was preserved). In order to enhance their motivation to participate, five times 10 Euro were disposed of by lot among the participants in Experiment 2 (irrespective of their additional participation in Experiment 3).

Procedure, Design, and Dependent Measures

In most procedural respects, Experiment 2 was identical to Experiment 1 except that it proceeded as a computer-controlled experiment for practical reasons and there was no filler task. The participants first read instructions equivalent to those in Experiment 1 and started to work on a computerized version of the bogus colour perception task. When finished, the computer program announced that it had calculated their score. To avoid having two versions of the computer program which would have to be counterbalanced across participants (in addition to the counterbalancing of matrix versions, see Experiment 1), however, they merely learned about the existence of two distinct perception categories ("colour sensitive" or "contrast sensitive") but not to which group they themselves belonged. Instead, the program explained that for the decision task to follow it was only necessary for them to know whether the persons to be rewarded belonged to their own group or to the outgroup. (This means at the same time that any possible identification with the ingroup should result from mere belongingness but not from any substantive features of the categories.) An explanation of the reward allocation task followed, using an example, and the participants also had the opportunity to go through the instructions for a second time if they wanted. After having finished the distribution matrices, the program reminded them that some of them would be asked to participate in a second session the next week. I delayed debriefing of all participants until this second session (Experiment 3) had been run. The whole procedure of Experiment 2 lasted about 45 minutes.

The standard condition of Experiment 2 - with the changes described above - was equivalent to the standard condition of Experiment 1. The irrelevance condition differed from the standard condition in only one respect: In the introduction of the reward allocation task, an added sentence stated that the group membership information given in the matrices would not be needed in the present session but only in the second and was retained here only for practical reasons. I highlighted this sentence in red to ensure that it would be noticed by the participants. All matrices and dependent measures were identical to Experiment 1.

Results & Discussion

As Table 2 shows, the amount of ingroup bias in the standard condition was comparable to Experiment 1, even though it reached only marginal significance because of the smaller sample size. Contrary to my expectations, ingroup bias did not vanish in the irrelevance condition but was even larger than in the standard condition. Thus, the results of Experiment 2 also failed to support the conversation logic approach to the minimal group paradigm.

Table 2: Average Ingroup Bias in Experiments 2 and 3 (Difference between Total Points Assigned to the Ingroup and Outgroup)

Experiment/Condition	Difference	SD	t	p[a]	d
Exp. 2 irrelevance (N = 28)	38.6	83.4	2.45	.01	.46
Exp. 2 standard (N = 24)	22.4	75.8	1.45	.08	.30
Experiment 3 (N = 24)	21.7	88.6	1.20	.12	.24

[a] One-tailed tests against the null hypothesis of no discrimination.

In retrospect, the most likely explanation for this might be the one also invoked in the discussion of Experiment 1: Although not intended, and indeed hoped to be circumvented by the new manipulation, the irrelevance manipulation in Experiment 2 might again have increased the salience of the group categorization, which in turn resulted in sizeable ingroup bias, over and above any possible reduction of it due to a perceived irrelevance of the group membership information for the allocation task.

Indeed, this suggests a fundamental difficulty in testing predictions of the conversation logic approach in the minimal group paradigm: It might not be possible at all to manipulate the perceived relevance of the group categorization without at the same time increasing its salience, because in order to manipulate the perceived relevance of the only piece of information that seems to be useful for the participants to guide their decisions, one must somehow relate to it, mention it, which might suffice to increase its salience and counteract the intended effect of the manipulation.

EXPERIMENT 3

Experiment 3 tested a unique prediction of the conversation logic approach, and one that should not be plagued with the problems discussed above. As outlined in the introduction, the conversation logic approach holds that the group membership information should induce the participants, by obeying the maxim of relevance, to make a difference in their reward allocations along the group categorization. However, it does not specify the direction of the difference, that is, pro-ingroup or pro-outgroup. Accordingly, it should be possible to systematically induce outgroup bias under suitable circumstances, at least in some participants. In Experiment 3, I tried to achieve this by creating a situation where the outgroup appeared more deserving of rewards than the ingroup, that is, a situation where a fairness or distributive justice norm is compatible with "making a difference". Consequently, participants acting according to such a norm (and typically, there are quite some participants in minimal group experiments found to distribute fairly; cf. e.g. Branthwaite, Doyle, & Lightbown, 1979; Turner, 1983) might be expected to exhibit outgroup bias.

Actually, this prediction was tested in the second session announced to the participants in Experiment 2. All participants in the standard condition of Experiment 2 were requested to take part in this second session, at the beginning of which they learned about the ostensible meantime result after the first session. They were told that up to this point the ingroup had been awarded about 25% more points than the outgroup. Because the participants were about to make allocation decisions in another round of distribution matrices, they had the opportunity to correct for this outgroup disadvantage by showing outgroup bias in their decisions if they wanted to. The latter should hold particularly for those participants who had distributed fairly in the first session.

Method

Participants

The 24 psychology undergraduates from the standard condition of Experiment 2 participated in what was for them the second session of their experiment (see method section of Experiment 2 for further details).

Procedure, Design, and Dependent Measures

The computerized instructions at the beginning of the session informed the participants that this second session was necessary because usually their concentration on this type of decision would decrease after about 30 matrices. They further learned that some participants had asked how many points both groups had received so far, and therefore we (the experimenters) had decided to announce the meantime result of both groups. Ostensibly, the ingroup had received 1809 points and the outgroup had received 1423 points. After this information, the experiment immediately proceeded with exactly the same set of matrices as in Experiment 2. Finally, after having finished the matrices, the participants received a post-experimental questionnaire similar to that used in Experiments 1 and 2. Together with the first session from Experiment 2, this additional session constituted a longitudinal design, with a major emphasis on changes in the participants' reward distributions.

Results & Discussion

Table 2 shows that the overall level of ingroup bias in Experiment 3 is largely unchanged from the first session (that is, the standard condition of Experiment 2). However, this first impression is not very informative with respect to the theoretical expectation entertained here, that is, that in particular those participants who had distributed fairly in the first session might exhibit outgroup bias in the second session.

More specific evidence relevant to this prediction can be gathered from a more refined analysis in terms of dominant distribution strategies of participants, as suggested in recent work (Blank, 2003; Petersen & Blank, 2001), which makes it possible to subdivide the sample in terms of dominant strategies in both sessions. The essence of this analysis (although the details are beyond the scope of this article) is, first, to identify the strategy with the largest pull score on a given matrix type. This is done on the basis of an expanded pull score analysis that includes a third pull score (in addition to the two conventional pull scores), which reflects the participant's tendency to check columns in the middle of the matrix (conversely, the two conventional pull scores reflect tendencies towards certain columns at the ends of the matrix). For example, the middle of a MIP & MJP vs. MD matrix represents the point of fairness, and a participant checking the middle column would therefore be assigned the maximum pull score for fairness.

The second important step in the analysis is to take the consistency of strategies across different matrix types in an experimental session into account (cf. the description of the matrix types in the method section of Experiment 1). Conceivably, if a strategy is dominant, it should be operating in all of the matrix types (usually, minimal group experiments make use of three different matrix types). Moreover, because each matrix type confounds two or more strategies by design, the cross-matrix type analysis helps to strip a dominant strategy from spurious companions, so to speak. In short, the dominant strategy analysis combines local dominance (i.e., within a given matrix type) and cross-matrix-type consistency to yield dominant distribution strategies of individuals at the level of an experimental session. In a validation study (Blank, 2003), such dominant strategies turned out to correspond quite well with the participants' self-reported strategies. However, it may also be the case that no dominant strategy is identified, as when participants respond randomly (Blank, 2003).

In Experiment 3, the dominant strategy analysis established that six of the 24 participants pursued a fairness strategy in the first session (including one participant who exhibited a mix between two cooperative dominant strategies, fairness and MJP). These participants are of main interest for the present purposes.[3] How did they behave in the second session? Two of them stuck to their fairness strategy, while the other four at least partially changed it in the predicted direction. More precisely, one participant completely shifted his or her strategy to a MOP (maximum outgroup payoff) strategy. This change - as identified on the basis of the objective reward allocations - was corroborated by the participant when asked about possible strategy changes in the post-experimental questionnaire: "In the second session, I tried to equalize the point scores of the groups and therefore always gave as much points as possible to the outgroup" [my translation]. The remaining three participants exhibited an inconsistent mixture between fairness and outgroup-favoring strategies (MOP or MDO - maximum differentiation in favor of the outgroup) in the second session. This partial change was also corroborated in the postexperimental questionnaire by one participant: "... in the second session occasionally more points to the outgroup, because it was behind in terms of the point score" [my translation]; the other two participants provided no relevant information.

Importantly, an outgroup-favoring strategy in the second session (MDO or MOP) was neither associated with any other consistent dominant session one strategy than fairness nor with any of the inconsistent strategy mixtures in session one. In other words, the changes predicted by the conversation logic account were in fact specific to the fair session one participants. This difference in outgroup favoritism proportions (four of six fair participants compared to none of the remaining 18 participants) is significant by a chi square test (corrected for small samples), chi square (1) = 10.00, $p = .002$. Thus, the results of the third experiment are more supportive of the conversation logic account than the results of the preceding two experiments, even though this support is not impressive in numbers and not all of the fair session one participants completely shifted to an outgroup-favoring strategy. However, it might be that some participants' desire to appear consistent across sessions had worked against the predicted changes and, therefore, the expectation of a complete and radical shift was too optimistic from the start. In sum, it seems fair to say that Experiment 3 yielded the first substantive support for my conversation logic analysis of the minimal group paradigm. Participants in the latter become inclined to differentiate in the first place, and when given a good reason to do so, they also differentiate in favor of the outgroup.

Seen from a slightly different angle: Fair participants discriminate if their underlying fairness motivation is compatible with making a difference. This may at the same time explain why they did not discriminate (in favor or against any of the two groups) in the first session: Their fairness motivation had suppressed any discriminatory demand that might have been conversation-logically conveyed. Once again, however, such a suppression mechanism would illustrate the comparative weakness of conversation logic effects in the MGP. They are easily overridden by the salience of the group membership information, and they seem to be just as easily suppressed by a fairness motivation under the standard MGP conditions.

GENERAL DISCUSSION

Let me summarize the main findings from the present experiments. Experiment 1 tested the conversation logic-based prediction that ingroup bias would be eliminated if the group categorization was not perceived as relevant for the reward allocation task. I tried to achieve this by making it explicitly relevant for a secondary task. As it turned out, however, the presence of a categorization-relevant secondary task heightened rather than diminished or eliminated ingroup bias. A second finding from Experiment 1 was that a cognitively demanding secondary task (whether categorization-relevant or not) reduced the amount of ingroup bias. In Experiment 2, the perceived relevance of the group membership information was manipulated without the help of a secondary task, by plainly telling the participants that this information was not relevant for it (but for a later task with the same materials). This new manipulation again led to more rather than less ingroup bias. My post hoc explanation for these unexpected findings was that any potential effect of perceived relevance of the group membership information for the reward allocation task was overridden by the increased salience of the group categorization. In retrospect, this unintended counter-effect seems to be an unavoidable consequence of the relevance manipulation.

However, with this interpretation I do not mean to immunize the conversation logic account against falsification. A counteracting salience effect, although possibly unavoidable in terms of experimental design, does not make it logically impossible for the participants to behave in accordance with the presumed conversation logic mechanisms (and in fact, the quote from one participant in Experiment 1 provided evidence that these mechanisms were possible to operate). Thus, the question is why the participants went on to use the group membership information for discrimination purposes even if they should, according to conversation logic, feel no need to do so. Two possibilities come to mind.

First, not all participants may in fact have perceived this reduced need. This may explain some of the Experiment 1 effects, since the induced relevance attributions to a secondary task did not logically exclude an attribution also to the primary task. Thus, some participants may have perceived the group membership information as relevant for both tasks. However, this explanation is less applicable to Experiment 2 because, in the irrelevance condition of this experiment, it was made quite explicit to the participants that the group membership information would not be needed for the matrix task.

Therefore, a second possibility seems more viable, namely, that at least some participants intentionally decided to use the group membership information in spite of its perceived irrelevance. Whatever the motivation behind such intentional decisions (I return to this issue below), their mere existence clearly indicates that the impact of conversation logic-based relevance perceptions is relatively weak in the minimal group paradigm, compared to other factors and processes.

On the positive side, Experiment 3 found support for a different prediction of the conversation logic account, namely that, if differentiation takes place, the direction of this differentiation is not restricted to ingroup favoritism but can also take on the form of outgroup favoritism under suitable circumstances. After having learned that the ingroup was "ahead" after the first session, participants who had distributed fairly in a first session shifted to outgroup bias in a second session. However, I should mention in all fairness that the conversation logic account cannot explain the whole pattern of results in Experiments on its own. Ironically, the very precondition for conversation logic-based outgroup favoritism to occur (i.e., fair distribution behavior in the first session, which means not differentiating) is left unexplained by it. That is, conversation logic mechanisms have to interact with other factors (as the impact of social norms like fairness) in order to produce the pattern of results in Experiment 3. While this does not invalidate the conversation logic account, it again testifies to its limited role in the minimal group paradigm.

Given that the conversation logic account can play, as we have seen, but a minor role in explaining the results of the complete set of experiments presented in this article, is there a better explanation? To begin with, social identity theory might well explain the participants' allocation behavior in Experiments 1 and 2, if we assume that the relevance manipulations had inadvertently increased the salience of the group categorization. This, in turn, would have led the participants to see themselves as group members and act accordingly, that is, exhibit intergroup discrimination. The fact that the participants showed less ingroup bias when they had to perform a cognitively demanding secondary task might also be interpreted in line with the social identity approach. It can be argued that this task offered them an opportunity to directly enhance their individual self-esteem by performing well, obviating the need for an indirect enhancement of self-esteem via identification with their minimal group. Consequently, it would be of no wonder that they showed no or less intergroup discrimination.

However, the assumption that a difficult secondary task would induce an individual self-esteem enhancement motivation in the participants is in itself clearly post hoc and cannot be verified by independent data in the present experiments. Moreover, even if this should have been the case, one might ask just why the participants found it more attractive to engage in interpersonal instead of intergroup behaviour. Or, why did the participants not try to pursue both personal and intergroup goals at the same time? Logically, this would have been entirely possible in this case. Finally, social identity theory cannot straightforwardly explain why some of the fair participants shifted towards outgroup bias in Experiment 3. I admit that this is not a big failure of social identity theory, because it never denied the impact of other than identity-enhancing motivations, like fairness, in the MGP. Then, granted the impact of fairness, the shift towards outgroup favoritism can simply be regarded as a situationally adapted form of fairness.

In general, however, what it is difficult to explain from the perspective of social identity theory is why there are such large differences in strategies between people, that is, why some individuals show ingroup favoritism, others distribute fairly or show outgroup favoritism, and still others pursue no meaningful intergroup strategy at all and allocate points randomly (cf. Blank, 2003). A similar explanatory problem exists if social norms are invoked to account for the participants' allocation behavior. In this case, one would have to explain why some people act according to a loyalty-to-the-ingroup norm, others according to a fairness norm, and so on.

Inherent Uncertainty of the Experimental Setting Leads to Strategy Variability

Perhaps the solution to this heterogeneity of behavior in research done with the MGP and the Tajfel matrices lies not in any "substantive" processes or mechanisms as suggested by social identity theory, social norm adherence, or conversation logic but in the inherent uncertainty of the experimental setting and the reward distribution task. Allocating points to people one does not even know, without any reasonable clue as to how to distribute besides the knowledge about those people's membership in one of two more or less meaningful categories, comes across for many participants as a rather strange and nonsensical task and creates considerable uncertainty as to the proper way to handle it. In fact, in the present studies, the question most often asked of the experimenters during the experimental sessions was "How am I to distribute the points? According to which criterion?". That is, there was clearly no self-evident "task solution" for many of the participants. In the face of such uncertainty, they may have sought to define the experimental situation in ways that (1) maximized sense and (2) minimized uncertainty.

As one way of meeting these criteria, they may have chosen to concentrate on the secondary memory task as an intuitively sensible task with a clear performance criterion ("remember as much and as correct as possible") and to more or less neglect the reward allocation task by responding randomly or according to some arbitrary criterion (e.g., always checking the middle column of the matrix). Alternatively, if concentrating on the reward allocation task, the participants might employ any simple strategy that makes sense within itself, that is, appears consistent and rational in the sense of conforming to some plausible and acceptable standard. Such standards may be social norms that are applicable to intergroup situations, like fairness or loyalty to the ingroup (see, e.g., Gaertner & Insko, 2001; Hertel & Kerr, 2001, on the impact of norms in minimal group situations), but also motivational standards like self-esteem enhancement, as suggested by social identity theory. Moreover, participants might use

conversationally implied situational cues or perceived demand characteristics to derive subjectively meaningful allocation strategies. Finally, personality differences may also play a role.

In short, when faced with an inherently uncertain situation, the participants look for and choose from an array of quite different cues and standards to guide their behavior, resulting in a multitude of distribution strategies. Of course, depending on the situation and on experimental manipulations, one or the other cue or standard may become influential, leading to mean shifts in strategies, without however reducing their variability.

Abrams and Hogg (1988) have presented a somewhat similar - at first glance - analysis of the minimal group situation (see also Grieve & Hogg, 1999; Hodson & Sorrentino, 2001; Jetten, Hogg, & Mullin, 2000). These authors, too, point to situational uncertainty as a major determinant of the participants' behaviour in the minimal group paradigm and hold that ingroup bias is a means of reducing the uncertainty of the relation between the two minimal groups. I agree with this analysis; however, I would extend it to the experimental situation as a whole, as described above.

That is, the first question is how the participants deal with this situation, how they define it in ways that maximize sense and minimize uncertainty. Such definitions may be in terms of individual, interpersonal or intergroup situations. Only if the participants come to define the situation as an intergroup situation arises the further question how to reduce its uncertainty in terms of intergroup strategies. Ingroup bias is one possibility, as conceived by Abrams and Hogg (1988). But this is not an inevitable consequence; fairness is another feasible strategy of dealing with it (and indeed, there were quite a few participants in the present experiments who chose fairness as their rationale for intergroup behavior rather than ingroup favoritism). Which solution the participants will endorse may depend on factors as conceived by social identity theory, for example, the degree of identification with the ingroup, but also on additional influences as self-presentation concerns. A participant may well identify with the ingroup and feel inclined to treat it more favorably but deliberately choose a fairness strategy because he or she assumes that the experiment has to do with ingroup bias and he or she does not want to appear prejudiced to the investigator.

In short, ingroup bias as a reaction to intergroup uncertainty is but one possible process in the minimal group paradigm which should be regarded within the larger context of the experimental situation and the participants' definition of it. A more detailed analysis, based on the post-experimental questionnaires, of these perceptions and the mechanisms that lead to one or another way of dealing with the experimental situation is currently under way. The conversation logic mechanisms featured in the present work are one possible mechanism in this process but, as we have seen, not a particularly powerful one. In any case, what follows from this analysis is that the minimal group paradigm, particularly when combined with the Tajfel matrices, is perhaps not the best way to study intergroup processes, because of its inherent uncertainty and the sometimes erratic behavior it provokes.

REFERENCES

- Abrams, D., & Hogg, M. A. 1988. Comments on the motivational status of self-esteem in social identity and intergroup discrimination. *European Journal of Social Psychology*, 18: 317-334.
- Billig, M., & Tajfel, H. 1973. Social categorisation and similarity in intergroup behaviour. *European Journal of Social Psychology*, 3: 27-52.
- Blank, H. 1997. Cooperative participants discriminate (not always): A logic of conversation approach to the minimal group paradigm. *Current Research in Social Psychology*, 2: 38-49.
- Blank, H. 1998. Memory states and memory tasks: An integrative framework for eyewitness memory and suggestibility. *Memory*, 6: 481-529.
- Blank, H. 2003. Identifying dominant allocation strategies of individuals in the minimal group paradigm. *Current Research in Social Psychology*, 9: 75-95.
- Bless, H., Strack, F. & Schwarz, N. 1993. The informative functions of research procedures: Bias and the logic of conversation. *European Journal of Social Psychology*, 23: 149-165.
- Bornstein, G., Crum, L., Wittenbraker, J., Harring, K., Insko, C. A., & Thibaut, J. 1983. On the measurement of social orientations in the minimal group paradigm. *European Journal of Social Psychology*, 13: 321-350.
- Bourhis, R. Y., Sachdev, I., & Gagnon, A. 1994. Intergroup research with the Tajfel matrices: Methodological notes. In M. P. Zanna & J. M. Olson (Eds.), *The psychology of prejudice: The Ontario Symposium* (Vol. 7, pp. 209-232). Hillsdale, NJ: Erlbaum.
- Branthwaite, A., Doyle, S., & Lighthown, N. 1979. The balance between fairness and discrimination. *European Journal of Social Psychology*, 9: 149-163.
- Brewer, M. B. 1979. In-group bias in the minimal group situation: A cognitive-motivational analysis. *Psychological Bulletin*, 86: 307-324.
- Diehl, M. 1990. The minimal group paradigm: Theoretical explanations and empirical findings. *European Review of Social Psychology*, 1: 263-292.
- Fiedler, K. 1988. The dependence of the conjunction fallacy on subtle linguistic factors. *Psychological Research*, 50: 123-129.
- Gaertner, L., & Insko, C. A. 2001. On the measurement of social orientations in the minimal group paradigm: Norms as moderators of the expression of ingroup bias. *European Journal of Social Psychology*, 31: 143-154.
- Gerard, H. B., & Hoyt, M. F. 1974. Distinctiveness of social categorization and attitude toward ingroup members. *Journal of Personality and Social Psychology*, 29: 836-842.

Grice, P. 1975. Logic and conversation. In P. Cole and J. L. Morgan (Eds.), *Syntax and semantics* 3: Speech acts (pp. 41-58). New York: Academic Press.

Grieve, P. G., & Hogg, M. A. 1999. Subjective uncertainty and intergroup discrimination in the minimal group situation. *Personality and Social Psychology Bulletin*, 25: 926-940.

Hertel, G., & Kerr, N. L. 2001. Priming in-group favoritism: The impact of normative scripts in the minimal group paradigm. *Journal of Experimental Social Psychology*, 37: 316-324.

Hertwig, R., & Gigerenzer, G. 1999. The "conjunction fallacy" revisited: How intelligent inferences look like reasoning errors. *Journal of Behavioral Decision Making*, 12: 275-305.

Hilton, D. J. 1995. The social context of reasoning: Conversational inference and rational judgment. *Psychological Bulletin*, 118: 248-271.

Hodson, G., & Sorrentino, R. M. 2001. Just who favors the in-group? Personality differences in reactions to uncertainty in the minimal group paradigm. *Group Dynamics: Theory, Research, and Practice*, 5: 92-101.

Jetten, J., Hogg, M. A., & Mullin, B.-A. 2000. In-group variability and motivation to reduce subjective uncertainty. *Group Dynamics: Theory, Research, and Practice*, 4: 184-198.

Mullen, B., Brown, R., & Smith, C. 1992. Ingroup bias as a function of salience, relevance, and status: An integration. *European Journal of Social Psychology*, 22: 103-122.

Petersen, L.-E., & Blank, H. 2001. Reale Gruppen im Paradigma der minimalen Gruppen: Wirkt die Gruppensituation als Korrektiv oder Katalysator sozialer Diskriminierung? [Real groups in the minimal group paradigm: Does the group context work as corrective or catalyst of social discrimination]. *Zeitschrift für Experimentelle Psychologie*, 48: 304-318.

Rabbie, J. M., Schot, J. C., & Visser, L. 1989. Social identity theory: A conceptual and empirical critique from the perspective of a behavioural interaction model. *European Journal of Social Psychology*, 19: 171-202.

Schiffmann, R., & Wicklund, R. A. 1992. The minimal group paradigm and its minimal psychology: On equating social identity with arbitrary group membership. *Theory & Psychology*, 2: 29-50.

Schwarz, N. 1999. Self-reports: How the questions shape the answers. *American Psychologist*, 54: 93-105.

St. Claire, L., and Turner, J. C. 1982. The role of demand characteristics in the social categorization paradigm. *European Journal of Social Psychology*, 12: 307-314.

Tajfel, H. 1970. Experiments in intergroup discrimination. *Scientific American*, 223: 96-102.

Tajfel, H., & Turner, J. C. 1986. The social identity theory of intergroup behaviour. In S. Worchel & W. G. Austin (eds.), *Psychology of intergroup relations* (pp. 7-24). Chicago: Nelson-Hall.

Tajfel, H., Billig, M. G., Bundy, R. P., & Flament, C. 1971. Social categorization and intergroup behaviour. *European Journal of Social Psychology*, 1: 149-178.

Turner, J. C. 1978. Social categorization and social discrimination in the minimal group paradigm. In H. Tajfel (Ed.), *Differentiation between social groups* (pp. 101-140). London: Academic Press.

Turner, J. C. 1983. Some comments on ... 'the measurement of social orientations in the minimal group paradigm'. *European Journal of Social Psychology*, 13: 351-367.

Turner, J. C. 1985. Social categorization and the self concept: A social cognitive theory of group behavior. In E. J. Lawler (Ed.), *Advances in group processes: Theory and research* (Vol. 2, pp. 77-122). Greenwich, CT: JAI Press.

Turner, J. C., Hogg, M. A., Oakes, P. J., Reicher, S. D., & Wetherell, M. S. 1987. *Rediscovering the social group: A self-categorization theory*. Oxford, UK: Blackwell.

ENDNOTES

[1] In different sessions, Stefan Röttger, Gregor Weißflog and myself served as experimenters.

[2] Note that this constitutes a difference to the procedure employed in Blank (1997), where the participants were required to remember the points allocated to the persons. This, however, often led the participants to choose numbers in the distribution matrices that were easy to remember, a strategy that obviously interfered with their allocation behaviour. Therefore, I tried to avoid this in the present studies.

[3] The other participants largely fell into three categories. Four individuals pursued other meaningful intergroup strategies like MDI, another four followed a consistent response tendency towards the middle of the matrices (which can be meaningfully distinguished from fairness in the dominant strategy analysis; see Blank, 2003), and finally, ten participants did not display any consistent strategy at all.

AUTHOR'S NOTE

The research reported in this article was supported by grant BL 493/1-1 from the DFG (German Research Foundation). I am grateful to Lars Petersen for helpful suggestions on a draft.

AUTHOR'S BIOGRAPHY

Hartmut Blank works as an Oberassistent (lecturer/senior lecturer) of social psychology at the University of Leipzig, Germany. His current research interests include the hindsight bias, eyewitness suggestibility, cognitive balance, and intergroup processes. E-mail address: blank@rz.uni-leipzig.de.