

Cognitive Science 48 (2024) e13440
© 2024 Cognitive Science Society LLC.
ISSN: 1551-6709 online
DOI: 10.1111/cogs.13440

Valence-Dependent Implicit Action Generalization Among Group Members

Jiecheng Huangliang,^{a,b} Yinfeng Hu,^{a,b} Xutao Zheng,^a Zikai Xu,^{a,b}
Wenying Zhou,^{a,b} Jun Yin^{a,b}

^a*Department of Psychology, Ningbo University*

^b*Center of Group Behavior and Social Psychological Service, Ningbo University*

Received 29 November 2023; received in revised form 28 February 2024; accepted 25 March 2024

Abstract

People implicitly generalize the actions of known individuals in a social group to unknown members. However, actions have social goals and evaluative valences, and the extent to which actions with different valences (helpful and harmful) are implicitly generalized among group members remains unclear. We used computer animations to simulate social group actions, where helping and hindering actions were represented by aiding and obstructing another's climb up a hill. Study 1 found that helpful actions are implicitly expected to be shared among members of the same group but not among members of different groups, but no such effect was found for harmful actions. This suggests that helpful actions are more likely than harmful actions to be implicitly generalized to group members. This finding was replicated in Study 2 by increasing the group size from three to five. Study 3 found that the null effect for generalizing harmful actions among group members is not due to the difficulty of detecting action generalization, as both helpful and harmful actions are similarly generalized within particular individuals. Moreover, Study 4 demonstrated that weakening social group information resulted in the absence of implicit generalization for helpful actions, suggesting the specificity of group membership. Study 5 revealed that the generalization of helping actions occurred when actions were performed by multiple group members rather than being repeated by one group member, showing group-based inductive

The study was approved by the Research Ethics Board of the Department of Psychology at Ningbo University and was performed in accordance with the relevant guidelines and regulations. All data, analysis code, and research materials have been made publicly available via the Open Science Framework and can be accessed at https://osf.io/azbvc/?view_only=028805d3046b475c97b8320d32d825f8. The design and analysis of Studies 1 and 2 were not preregistered, but Studies 3–5 were preregistered and can be accessed at https://osf.io/34z6j?view_only=76aec648c6e84ab3909b694a6955f4f4. Notably, Studies 3a and 3b in the current main text refer to Studies 1a and 1b in the preregistered studies, and Studies 4 and 5 in the current main text refer to Studies 2 and 3 in the preregistered studies.

Correspondence should be sent to Jun Yin, Department of Psychology, Ningbo University, No. 616 Fenghua Rd., Ningbo 315211, China. E-mail: yinjun1@nbu.edu.cn

generalization. Overall, these findings support valence-dependent implicit action generalization among group members. This implies that people may possess different knowledge regarding valenced actions on category-based generalization.

Keywords: Action generalization; Helpful action; Harmful action; Social group; Implicit generalization

1. Introduction

Humans are profoundly dependent on social groups for affiliation with others and to collectively act. This living environment has evolutionarily shaped the cognitive capacities to process and represent social groups (Hirschfeld, 2001). Social groups can automatically activate categorical thinking to guide our interpretations and expectations of the actions of others (Bodenhausen, Kang, & Peery, 2012; Macrae & Bodenhausen, 2000). People believe that individuals within a social group share behavioral tendencies, and both adults and infants implicitly generalize the actions and attributes of known individuals in a social group to unknown members (Liberman, Woodward, & Kinzler, 2017; Powell & Spelke, 2013; Ranganath & Nosek, 2008; Yin, Duan, Huangliang, Hu, & Zhang, 2022). This type of inductive generalization, which is a form of category-based inference, allows for efficient learning about the social world and the ability to interact appropriately in social situations by acquiring the social group to which individuals belong, even without specific knowledge of each individual (Liberman et al., 2017). However, as group members have their own desires for executing actions, sometimes judging people based on the actions of group members with whom they are associated is not guaranteed to generate accurate evaluations (Banaji & Bhaskar, 2000; Navon, Shechter, & Bar-Anan, 2021). It is thus important to understand the boundary conditions that enable us to perform inductive generalization for social group members (Crawford, Sherman, & Hamilton, 2002; Duan et al., 2021), especially from an implicit process perspective, as it has been suggested that implicit processes reflect our automatic responses and guide our daily behaviors toward processed targets more than their explicit counterparts (Frith & Frith, 2008; Greenwald, Poehlman, Uhlmann, & Banaji, 2009; Kurdi, Mann, Charlesworth, & Banaji, 2019).

1.1. Generalizing actions among group members

Social groups are constructed as social kinds or categories, providing the basis for category-based inference (Billig & Tajfel, 1973; Bodenhausen et al., 2012; Liberman et al., 2017). Specifically, people use their knowledge of social category membership to inductively generalize attributes to novel members. This traditional field within social psychology has generated numerous studies documenting the formation of stereotypes associated with social groups and their application to unknown members without considering individual characteristics (Chan et al., 2012; Greenwald & Banaji, 1995; Jussim, Crawford, & Rubinstein, 2015). Even the immediately learned attributes of an individual, including attitudes and traits, are readily generalized to other group members (Chen & Ratliff, 2015; Crawford et al., 2002; Ranganath & Nosek, 2008; Ratliff & Nosek, 2011). However, studies describing stereotypes

have placed more emphasis on the stable attributes or traits underlying behaviors being extended within or between group members (Kruse & Degner, 2021). However, we usually predict and expect others' actions in transient settings, in which understanding actions may appeal not to stable traits but to temporary situational demands (e.g., group similarities and norms) and situation-specific short-term states (e.g., goals and intentions) (Kruse & Degner, 2021; Westra, 2019). Recent research has directly shown that social groups guide our expectations of group members' actions in transient settings, and group members are implicitly expected to act alike, showing implicit action generalization (Duan et al., 2021; Powell & Spelke, 2013; Xu, Duan, Chen, & Yin, 2019; Yin et al., 2022). Duan et al. (2021) described two social groups of agents, where two agents from the same group performed a similar action (i.e., jumping or going straight), and the remaining agent of this group behaved either consistently or inconsistently with the others. Group-consistent actions were identified more quickly than group-inconsistent actions. Thus, a new social group member was implicitly expected to engage in actions similar to those of known members; otherwise, the identification of the unknown agent's actions should be independent of the actions of the known group members. Such action generalization has been found very early in development (Powell & Spelke, 2013) and, importantly, was not observed when social group cues were removed (Duan et al., 2021; Xu et al., 2019). Moreover, further studies have revealed that implicit action generalization is based on common goals among group members, rather than shallow-level movements (Yin et al., 2022, 2023; Hu et al., 2023).

However, the actions mentioned above were aimed at individual (i.e., nonsocial object-directed) goals without any direct impact on others. Actions in reality can have social goals that result in beneficial or detrimental consequences for others and can be evaluated with different valences (Gilbert & Basran, 2019; Jara-Ettinger, Gweon, Schulz, & Tenenbaum, 2016). The evaluation of social action as helpful or harmful is crucial for human biological adaptation and can be linked to positive or negative valence in moral thought (Haidt & Kesebir, 2010; Schein & Gray, 2018). However, it remains unclear whether action valence influences implicit generalizations for social actions. In other words, it is unknown whether implicit generalizations for social actions are primarily driven by accumulated observations of shared attributes (e.g., similarity), regardless of the action's valence. Addressing this question can clarify whether individuals possess distinct knowledge regarding valenced actions in the context of action generalization and reveal whether a specific social action holds a privileged status as an attribute shared among group members.

1.2. Action generalization is possibly valence-dependent

Intuitively, the generalization of helpful and harmful actions should be similarly observed, showing valence-independent action generalization, as people who share the same behavioral tendency are often clustered together, and both actions may be treated as common attributes of social groups (Brewer, Hong, & Li, 2004), for example, "good" and "bad" groups. This speculation largely emphasizes the possible similarities within a social group, that is, the overall similarity between the individual who serves as the basis for the generalization and the individual who serves as the potential target. However, in terms of category-based inference,

people can mostly access their knowledge of properties. Thus, inductive generalization is guided not simply by overall similarity but also by people's knowledge; for example, knowledge of causal connections and properties plays a crucial role in guiding the generalization of information (Bright & Feeney, 2014; Heit & Rubinstein, 1994; Kemp & Tenenbaum, 2009). In this case, if individuals possess distinct knowledge of actions with varying valences, such actions may be selectively generalized among group members.

Positive and negative information are suggested to be processed differently and construed as two asymmetrical classes in the mind (Alves, Koch, & Unkelbach, 2017b; Unkelbach, Koch, & Alves, 2019). Generally, negative information is thought to have a stronger psychological impact than positive information (Alves et al., 2017b; Baumeister, Bratslavsky, Finkenauer, & Vohs, 2001). For instance, negative social information receives more attention than positive social information (Noyes & Keil, 2018; Pratto & John, 1991), and stimuli with negative valence are more accurately memorized (Robinson-Riegler & Winton, 1996). This negativity bias has recently been explained by the finding that people have different causal knowledge for positive and negative actions, in which a negative action can provide specific and diagnostic information for the individual who performs it (Alves, Koch, & Unkelbach, 2018; Fazio, Pietri, Rocklage, & Shook, 2015; Taylor, 1991; Unkelbach, Alves, & Koch, 2020). Specifically, Alves and colleagues (2017b) found that negative attributes are less prevalent than positive attributes in the social world since people generally conform to social norms that promote positive behaviors, making negative, norm-incongruent behaviors less common. People's mental representations of the social world are suggested to mirror this ecological structure information of positivity prevalence. In this case, positive attributes should have a greater probability of being shared among relevant persons than negative attributes, but negative information has a greater impact and is diagnostic in processing a particular individual, as it is diverse and distinct (Alves et al., 2018; Unkelbach et al., 2020).

This type of knowledge may be integrated with category-based inference, guiding action generalization. Accordingly, the negative action could be largely bound to the particular individual, showing less or even no group-based generalization for negative actions. The combination of this knowledge and category-based inference implies that helpful actions, as positive actions, would be more likely to be generalized than harmful actions, as negative actions, across social groups. Therefore, we formulated the hypothesis that implicit generalizations for social actions are not solely driven by accumulated evidence of shared attributes. Instead, we propose that these generalizations are also influenced by the knowledge of valence asymmetries, as discussed above. Given that knowledge of valence asymmetries is derived from both ecological structure information and normative knowledge (Alves et al., 2018; Unkelbach et al., 2020), this characteristic may unconsciously shape our evaluations of the social world. Accordingly, we predicted that helpful actions would be more likely than harmful actions to be implicitly generalized within social groups, thus demonstrating valence-dependent implicit action generalization.

1.3. Overview of the current research

To explore whether helpful and harmful actions are implicitly expected to be shared within social groups to a similar degree, the present research employed visual scenes featuring

computer-animated events to simulate social groups and their actions, which was inspired by previous research (Duan et al., 2021; Hamlin, Wynn, & Bloom, 2007; Ludwig et al., 2020; Powell & Spelke, 2013). Helpful (harmful) actions were presented as helping (hindering) another as previously developed to explore the social evaluations of individuals based on their social interactions (Hamlin et al., 2007; Isik, Koldewyn, Beeler, & Kanwisher, 2017). There is extensive documentation indicating that people can comprehend social meaning in these animations (Duan et al., 2021; Isik et al., 2017). In the present research, a cartoon animation was used to simulate social groups and types of actions. Visual cues, namely, spatial proximity and synchronous movement, which are considered key characteristics governing groupness (Campbell, 1958), were used to create social groups (Duan et al., 2021; Powell & Spelke, 2013; Yin et al., 2022). After introducing social groups, two members from a social group then successively performed an action to help (hinder) an agent who does not belong to a group climb a hill. Later, a new member moved out and paused at the foot of the hill, which appeared to consider whether to help (hinder), and implicit action generalization was tested.

In the current research, our focus was on participants' spontaneous and implicit responses when generalizing helpful and harmful actions displayed by group members. Explicit measurements¹ often rely on consensual cultural stereotypes and norms (Kurdi et al., 2019), while implicit methods tend to capture processes that dominate our automatic responses to processed targets beyond what is assessed with explicit measures (Frith & Frith, 2008; Greenwald et al., 2009; Kurdi et al., 2019). We used the lexical decision task as an indirect measure to assess implicit action generalization. Specifically, during the test, a pseudoword or a real word naming a helping, hindering, or neutral interaction was presented. The implicit generalization of helpful and harmful actions across group members was measured by asking the participants to judge whether this word was real or not instead of explicitly asking what would be shared by group members. The rationale is that if participants implicitly expect the target to perform the same action as the exemplars, the real word with the helping or hindering meaning would be identified more quickly than the neutral word. This facilitation effect is expected because the meaning of the helping or hindering words corresponds to the content of the implicitly induced generalization, similar to semantic priming. This rationale of implicit measurement is commonly used in studies that involve processing emotions using pictures and emotion-related words (Castner et al., 2007; Hill & Kemp-Wheeler, 1989; Lorenz & Newman, 2002).

Five studies were included in this research. First, Study 1 conducted two separate experiments to assess whether helpful and harmful actions were implicitly generalized to a new group member in a similar manner. In Study 1a, we measured implicit action generalization for helpful actions via the facilitation effect in the lexical decision task with real words of helping and neutral interaction. In Study 1b, we measured implicit action generalization for harmful actions via the same task but with real words of hindering and neutral interaction. Then, Study 2 attempted to replicate the results of Study 1 by increasing the group size, while Study 3 was designed to rule out possible alternatives (such as the paradigm's ability to detect the expectation for hindering action) by comparing between-person and within-person action generalization. Finally, we explored whether the implicit action generalization within social

groups was specifically dependent upon group membership (Study 4) and whether it relied on group- rather than individual-based inference (Study 5).

2. Study 1

This study tested whether the helpful or harmful actions of known group members would be implicitly generalized to an unknown member. Valence-dependent implicit action generalization predicted that the facilitation effect should be observed for helping actions, but no such effect or a weaker effect would be found for hindering actions. Such a prediction should be observed only in group members with whom the target was associated, as inductive generalization relies on social category membership. Here, we tested the implicit generalization for helping and hindering actions with different subjects instead of a within-subject design, as within-subject comparisons for the two types of actions may highlight the valence differences between them and make the helpful or harmful information salient.

2.1. Study 1a

2.1.1. Methods

2.1.1.1. Transparency and openness: We report how we determined our sample size, all data exclusions (if any), all manipulations, and all measures in the study, and we follow Journal Article Reporting Standards (JARS; Kazak, 2018). All data, analysis codes, and research materials were made publicly available via the Open Science Framework and can be accessed at https://osf.io/azbvc/?view_only=028805d3046b475c97b8320d32d825f8. The data were analyzed using Python, version 3.9.7 (Van Rossum & Drake, 2009), with the package Pingouin, version 0.5.0 (Vallat, 2018) and the package Plotnine, version 0.8.0 (<https://github.com/has2k1/plotnine>). The design and analysis of Studies 1 and 2 were not preregistered, but Studies 3–5 were preregistered and can be accessed at https://osf.io/34z6j?view_only=76aec648c6e84ab3909b694a6955f4f4. Notably, Studies 3a and 3b in the current main text refer to Studies 1a and 1b in the preregistered studies, and Studies 4 and 5 in the current main text refer to Studies 2 and 3 in the preregistered studies.

2.1.1.2. Participants: Thirty adult volunteers (12 men and 18 women, mean age = 20 years, aged between 18 and 25 years) were paid approximately 20 RMB each to participate in this study. All participants were of Han ethnicity, the majority ethnic group in China. The sample size was determined as follows. We ran a pilot experiment that was similar to the current experiment, except that only members of the same group as the two exemplars were required to be judged and their actions expected (please see Study S1 in the Supplementary Material for details), and 20 participants participated in this pilot study. The observed F value for the difference between the helping action word and the neutral action word conditions was 14.31. We used the BUCSS software package (Anderson, Kelley, & Maxwell, 2017) to perform the power analysis for the within-subject ANOVA in our study

(https://designingexperiments.shinyapps.io/BUCSS_ss_power_wa/). Specifically, we set the observed F value to 14.31 and focused on the 2×2 interaction effect between the agent group and word meaning. The target alpha level for the previous and planned studies was 0.05. The desired statistical power was set to 0.8, with a level of assurance of 0.8, which means that we had an 80% chance of obtaining the desired statistical power or higher with the estimated sample size. With the above settings, the calculated sample size was 30. Considering that all studies had the same repetitive measurements as the current study, each study recruited the same number of participants as the current study. That is, the sample size in each study was 30.

All participants had normal or corrected-to-normal vision and no history of neurological disorders. The study was approved by the Research Ethics Board of the Department of Psychology at Ningbo University and was performed in accordance with the relevant guidelines and regulations.

2.1.1.3. Apparatus and stimuli: The stimuli consisted of cartoonized animations depicting helping events and were made in Blender (version 3.0.1). Each event had a background containing two hills, one at the left edge and the other at the right edge of the screen. Positioned behind these hills, in the center of the screen, were two groups, each consisting of three agents acting as potential providers of help (i.e., actors). These actors shared the same appearance but differed between the two groups. Moreover, each event involved three additional distinct agents positioned at the bottom of the screen, who served as patients in need of help. To create two groups of agents and match them across conditions, 10 pairs of agents were generated from a pool of five predefined agents with different appearances (i.e., blue-green square, blue cross, pink octagon, yellow star, and lilac decagon). For the trials in each condition, a pair of agents was randomly sampled from these 10 pairs, and the remaining three distinct agents were set as the patients. This randomized sequence remained the same in each condition to match possible physical differences. Each agent was rendered with two eyes as facial features and subtending approximately $1.0^\circ \times 1.0^\circ$.

The stimuli were all presented on a gray background using a 19-inch cathode ray tube monitor (1024×768 pixel resolution; 85 Hz refresh rate) at a 60-cm viewing distance using custom software written in MATLAB with Psychophysics Toolbox libraries (Brainard, 1997).

2.1.1.4. Procedure and design: Each trial started with the presentation of a fixation cross ($1.5^\circ \times 1.5^\circ$) in the center of the screen for 300–400 ms. Subsequently, two groups of agents were displayed, one on each side at the top of the screen, while an additional three agents, acting as patients, were positioned at the bottom (foreground) of the screen. Initially, the agents within each group were in close proximity but exhibited greater spatial distance between the two groups. Following this, the agents in each group immediately executed a series of synchronized movements along a circular path lasting approximately 1450 ms and then returned to their original positions atop the platform. This synchronized movement served as a visual cue to the observers, indicating that they were a “group” (refer to Fig. 1, “introducing social groups,” and Videos S1 and S2 in the Supplementary Material).

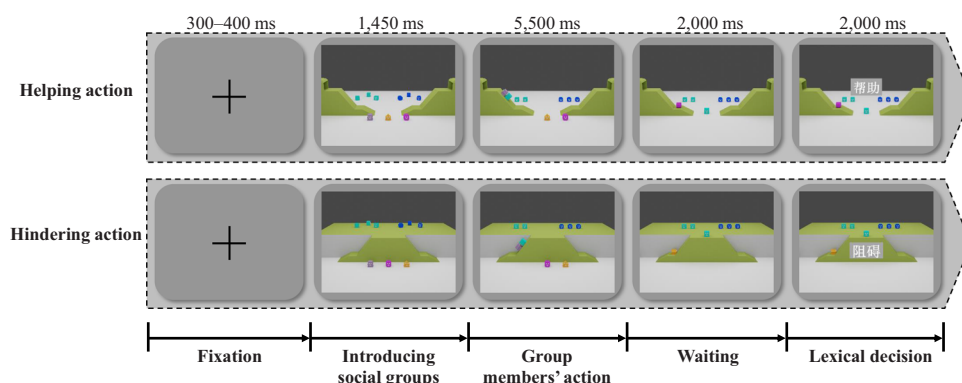


Fig. 1. An example trial sequence for helping and hindering actions. The trial begins with a fixation and displays the movements to introduce the social groups (“introducing social groups” phase). Then, in the “group members’ action” phase, two agents from the social group successively perform the same helping or hindering action. Finally, after waiting a few seconds, a Chinese word with two characters is presented at the center top of the platform, and the participants are asked to judge whether this word is real or pseudo by pressing the appropriate button (“lexical decision” phase).

Following this introductory sequence, one of the three patients moved to the bottom of a hill, either at the left or right (balanced within participants), and attempted to reach the top. Despite repeated attempts, the patient slid partway down the steep slope and could not reach the summit successfully. During this time, one of the actors initially moved toward the bottom of the hill at the center of the screen and helped the patient reach the hill’s summit. This patient then moved into the hole at the top of the hill and vanished, while the actor returned to its starting point by following the original path. Subsequently, another patient and a new actor from the same group as the first actor reenacted similar helping events. This phase lasted approximately 5500 ms (see Fig. 1; “group members’ actions”). Next, a third patient tried to climb the hill, mirroring the actions of the previous two patients. A new actor, who was from either the same or a different group as previous two actors (i.e., known members), went to the bottom of the hill at the center of screen, in a position where they could help, and paused. There was a pause of 2000 ms (refer to Fig. 1; “Waiting”), during which no explicit task was assigned. The aim of this phase was to induce expectations for the agent’s future action. Finally, a Chinese word with two characters was displayed at the center top of the platform. The participants were asked to determine whether this word was real or pseudo by pressing the “J” or “F” button, respectively, on a standard keyboard (see Fig. 1; “lexical decision”). If the participant did not respond within 2000 ms, the program automatically proceeded to the next trial. The response time (RT) and accuracy were recorded.

The participants completed 104 trials. Each agent group condition (same vs. different) had 40 trials presenting real words and 12 trials presenting pseudowords. For the trials presenting real words, half of them indicated an “interaction with helpful outcomes” (i.e., helping action words; e.g., “帮助” “bangzhu” means “help”), and half of them indicated an “interaction

with neutral outcomes” (i.e., neutral action words; e.g., “交谈” “jiaotan” means “converse”). For each word type, five words were collected in a pilot study, and the familiarity, word frequency, and number of strokes were matched between the helping action and the neutral action words (please see Study S2 in the Supplementary Material). The pseudowords were created by combining two characters from two words, which were sampled from five helping and five neutral action words. Finally, four pseudowords were selected by two experts in Chinese. Each real word was repeated four times in each agent group condition, maintaining a balance across conditions. The trials presenting pseudowords randomly sampled a word from the pool of four pseudowords but were constrained so that each word was presented an equal number of times. The trials presenting pseudowords were used as fillers and supported the word judgment task, and they were not related to our measurement. In this case, when only the trials presenting real words were considered, a 2 (agent group: same vs. different) $\times$ 2 (word meaning: helping action vs. neutral action) within-subject design was obtained, and each condition had 20 trials.

To reduce the total time of the experiment but to illustrate that the last actor will indeed have a future action, only 20% of the trials of each condition displayed the outcome of whether the last actor helped the patient. Without this display, there was a risk of participants forming the expectation that the last actor only moved to the middle of two hills. This action was shown after participants completed the lexical decision task. The trials of whether the last actor provided help were equally distributed. When the last actor did not help the patient, the actor went to another hill that no one was climbing. The Supplementary Videos showed the trials in which the outcome was not present. Before the formal experiment, the participants had 2 min to practice, during which the trials were randomly sampled from the formal trials. Each participant was given a short break after completing 26 trials and then completed the experiment within 30 min.

2.1.1.5. Analysis: As described above, 24 trials presenting pseudowords were used as fillers and were not the main focus, and 80 trials (20 trials in each condition) presenting real words were used as implicit measurements to detect the participants’ action generalizations. Since the participants’ reaction times (RTs) were emphasized and closely related to our hypothesis, the analyses in all experiments focused primarily on the RT data after excluding the trials with incorrect responses, as was previously done (Duan et al., 2021; Yin et al., 2022); the accuracy data are also reported in the Supplementary Material.

Moreover, given that the trials were nested within the participants, mixed-effects models that included trial-level data rather than ANOVAs on summary scores were used and yielded results similar to those reported by the ANOVAs (please see the results of using mixed-effects models on RT data in the Supplementary Material). In addition, we computed the Bayes factor (BF_{10} , $H1/H0$ as computed here) using JASP software (JASP Team, 2021; Version 0.16) for critical comparisons between helping/hindering and neutral action words. According to the Bayesian analysis, $BF_{10} < 0.33$ suggests substantial evidence against between-condition differences, while $BF_{10} > 3.00$ suggests substantial evidence supporting between-condition differences. Values between 0.33 and 3.00 are considered inconclusive (Dienes, 2014).

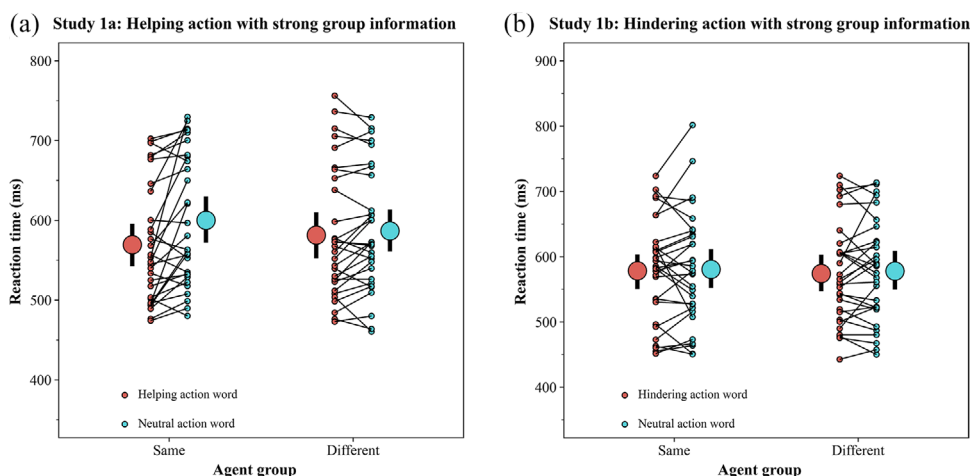


Fig. 2. Mean response time (RT) for the correct responses in (a) Study 1a and (b) Study 1b as a function of the agent group and word meaning. The dots indicate the mean RT for each participant in each condition and are connected for the paired conditions. The large circles denote the means, and the black vertical lines indicate the CIs of the means.

2.1.2. Results and discussion

The mean RT of each participant for each condition was computed after excluding the trials with incorrect responses (exclusion rate = 2.9%) and is shown in Fig. 2a. The RT data were subjected to a repeated-measures 2 (agent group: same vs. different) $\times$ 2 (word meaning: helping action vs. neutral action) ANOVA. The main effect of the agent group was not significant ($F(1, 29) = 0.01, p = .916, \eta_p^2 < 0.01$), but the main effect of the word meaning was ($F(1, 29) = 13.23, p = .001, \eta_p^2 = 0.31$), showing that the participants responded more rapidly in the helping action word condition ($M = 575$ ms, $SE = 14$) than in the neutral action word condition ($M = 593$ ms, $SE = 14$; 95% confidence interval [CI] of the mean difference = [9, 28]). More importantly, the interaction effect was significant ($F(1, 29) = 6.72, p = .015, \eta_p^2 = 0.19$). The simple effects analysis revealed that in the same group condition, the participants responded more quickly when the meaning of the words was related to a helpful interaction ($M = 569$ ms, $SE = 14$) than when it was related to a neutral interaction ($M = 600$ ms, $SE = 15$; $t(29) = 3.50, p = .002$, Cohen's $d = 0.64$, 95% CI of the mean difference = [-49, -13]; i.e., facilitation effect), showing a facilitation effect. However, in the different group conditions, there was no significant difference in the RTs for helping ($M = 581$ ms, $SE = 15$) and neutral action words ($M = 587$ ms, $SE = 14$; $t(29) = 1.19, p = .246$, Cohen's $d = 0.22$, 95% CI of the mean difference = [-14, 4]). The Bayesian factors measuring the difference in the RTs for the helping and neutral action words were 23.10 (0.37) when the agent was a member of the same (a different) group as the two known individuals. Hence, we established that helpful actions are implicitly expected to be shared within social groups. The null effect in the different group condition also ruled out the possibility that helping-related words were always identified as real more quickly than neutral words.

2.2. Study 1b

2.2.1. Methods

Thirty new adult volunteers (14 men and 16 women, mean age = 19 years, aged between 17 and 25 years) were paid approximately 20 RMB each to participate in this study. All participants were of Han ethnicity, the majority ethnic group in China. The procedure and design were the same as in Study 1a, except that a hindering event was presented, and the five helping action words were replaced by five words indicating an “interaction with hindering outcomes” (i.e., hindering action words; e.g., “阻碍,” “zuai” means “hinder”). In hindering events, two groups of agents were positioned on a platform atop hills, while another three agents successively attempted to climb a hill and reach the summit of the platform. However, the two actors from the same group on the platform hindered the patients’ attempts to reach the platform and blocked their paths to the platform (see Fig. 1 and Videos S3 and S4). The blocked patient descended from the hill and exited the screen, while the actor returned to his original position.

To test whether the helping and hindering actions were perceived as having different valences, 30 new participants (13 men and 17 women, mean age = 19 years, aged between 17 and 22 years) were recruited to watch two videos: an agent helping or hindering another agent climbing a hill. The participants then answered the question, “How do you evaluate the actor’s actions?” (rated from -3 = strongly negative to 3 = strongly positive). The presentation order of the two videos was counterbalanced across the participants. A paired t test revealed that the valence of the helping action ($M = 2.30$, $SE = 0.17$) was more positive than that of the hindering action ($M = -2.07$, $SE = 0.24$; $t(29) = 14.49$, $p < .001$, Cohen’s $d = 2.65$, 95% CI of the mean difference = $[3.75, 4.98]$). The valence of the helping action was greater than zero ($t(29) = 13.76$, $p < .001$, Cohen’s $d = 2.51$), the valence of the hindering action was less than zero ($t(29) = 8.81$, $p < .001$, Cohen’s $d = 1.61$), and the intensity (i.e., absolute value of the rated valence) of the helping action was not significantly different from that of the hindering action ($t(29) = 0.41$, $p = .682$, Cohen’s $d = 0.08$). These results suggest that the manipulation of the action valence in our studies is valid and that the intensity matches.

2.2.2. Results and discussion

2.2.2.1. Within-experiment tests: Fig. 2b shows the mean RTs for the trials with correct responses in all conditions. The RT data were submitted to a repeated-measures 2 (agent group: same vs. different) $\times$ 2 (word meaning: hindering action vs. neutral action) ANOVA after removing the trials with incorrect responses (exclusion rate = 1.8%). No effects were significant (agent group: $F(1, 29) = 0.47$, $p = .499$, $\eta_p^2 = 0.02$; word meaning: $F(1, 29) = 0.26$, $p = .617$, $\eta_p^2 < 0.01$; interaction effect: $F(1, 29) = 0.06$, $p = .809$, $\eta_p^2 < 0.01$). The Bayesian factors measuring the difference in the RTs between the hindering and neutral action words were 0.20 (0.23) when the agent was from the same (a different) social group. The equivalence tests for the RT difference between the harmful action word and neutral action in both agent group conditions suggested statistical equivalence (please see the results of the equivalence tests in the Supplementary Material).

2.2.2.2. Between-experiment tests: To further test whether action valence influences implicit action generalization, a 2 (action valence as a between-subject factor: helpful in Study 1a vs. harmful in Study 1b) $\times$ 2 (agent group as a within-subject factor: same vs. different) $\times$ 2 (word meaning as a within-subject factor: valenced action vs. neutral action) ANOVA was conducted for the RT data, and a three-way interaction was found ($F(1, 58) = 4.90, p = .031, \eta_p^2 = 0.08$). This interaction was due to a two-way interaction between the agent group and word meaning when the helpful action was present in Study 1a but not when the harmful action was present in Study 1b. Moreover, we analyzed the three-way interaction by fixing the variable for the agent group and found that the interaction between the action valence and word meaning was significant when the expected agent was from the same group as the two exemplars ($F(1, 58) = 6.17, p = .016, \eta_p^2 = 0.10$) but not when this agent was from a different group ($F(1, 58) = 0.03, p = .859, \eta_p^2 < 0.01$). To further analyze this significant two-way interaction effect, we calculated the facilitation effect by subtracting the mean RT of identifying words indicating valenced interactions from that of identifying neutral words. Our findings revealed that when the expected agent was from the same group as the two exemplars, the facilitation effect of the helpful action in Study 1a ($M = 31$ ms; $SE = 9$) was greater than that of the harmful action in Study 1b ($M = 2$ ms; $SE = 7$; $t(58) = 2.48, p = .016$, Cohen's $d = 0.64$, 95% CI of the mean difference = [5, 52]).

Hence, under conditions similar to those in which the participants successfully formed an implicit generalization for a helpful action, the participants tended to not implicitly generalize a harmful action within social groups. This result does not mean that the harmful action cannot always be implicitly generalized. Rather, it suggests that the action valence modulates the implicit action generalization, and the implicit generalization for the harmful action does not come as easily as that for helpful action, at least in the current dynamic scenes, supporting valence-dependent implicit action generalization.

3. Study 2

Study 1 revealed valence-dependent implicit action generalization among group members. However, the event depicted very small groups, and the two sampled agents were the majority. Hence, when participants observed a third group member coming forward to act, the perceived obligation of group members to perform helping actions may have been more pronounced, and the valence-dependent implicit action generalization may be specific to this setting. This study aimed to replicate previous findings by having groups of five members, and participants were tasked with forming expectations about a third group member based on a sample of only two members, which was a minority in the group.

3.1. Study 2a

3.1.1. Methods

Thirty new adult volunteers (15 men and 15 women, mean age = 21 years, aged between 18 and 26 years) were paid approximately 20 RMB each to participate in this study. All participants were of Han ethnicity, the majority ethnic group in China. The procedure and design

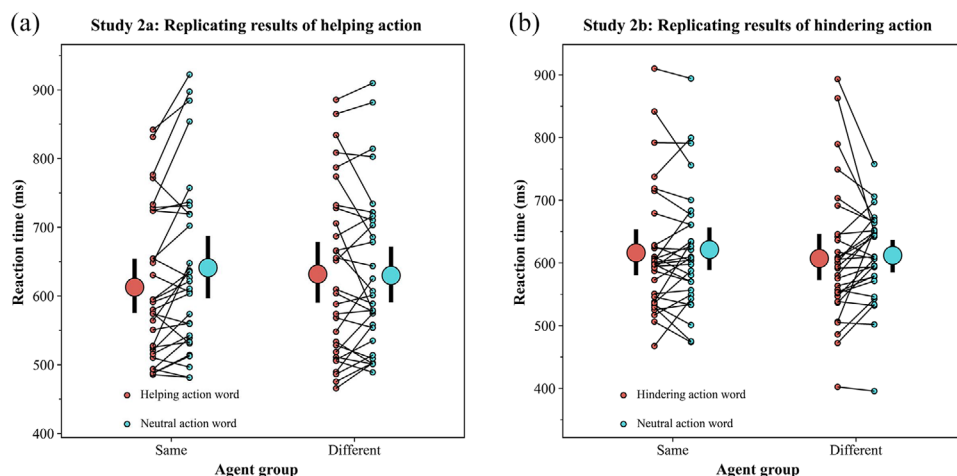


Fig. 3. Mean response time (RT) for the correct responses in (a) Study 2a and (b) Study 2b as a function of the agent group and word meaning. The dots indicate the mean RT for each participant in each condition and are connected for the paired conditions. The large circles denote the means, and the black vertical lines indicate the CIs of the means.

were similar to those of Study 1a, except that each group had five agents, and participants were involved in forming expectations about a third agent from either the same group (Video S5) or a different group (Video S6) based on a sample of two other group members.

3.1.2. Results and discussion

The mean RT for each participant in each condition was calculated, excluding the trials with incorrect responses (exclusion rate = 3.9%), and the results are shown in Fig. 3a. The RT data were analyzed using a repeated-measures 2 (agent group: same vs. different) $\times$ 2 (word meaning: helping action vs. neutral action) ANOVA. The main effects of agent group ($F(1, 29) = 0.83, p = .370, \eta_p^2 = 0.03$) and word meaning ($F(1, 29) = 3.41, p = .075, \eta_p^2 = 0.11$) were not significant. However, the interaction effect was significant ($F(1, 29) = 9.35, p = .005, \eta_p^2 = 0.24$). Further analysis revealed that in the same group condition, participants responded faster when the word meaning was related to a helpful interaction ($M = 613$ ms, $SE = 20$) than when it was related to a neutral interaction ($M = 641$ ms, $SE = 23$; $t(29) = 3.23, p = .003$, Cohen's $d = 0.59$, 95% CI of the mean difference = $[-46, -10]$). In contrast, there was no significant difference in the RTs for the helping ($M = 632$ ms, $SE = 23$) and neutral action words ($M = 630$ ms, $SE = 21$) in the different group condition ($t(29) = 0.23, p = .823$, Cohen's $d = 0.04$, 95% CI of the mean difference = $[-16, 19]$). Bayesian factors measuring the difference in the RTs for the helping and neutral action words were 12.28 (0.20) when the agent was a member of the same (a different) group as the two exemplars. These results suggest that helping actions are implicitly expected to be shared within social groups.

3.2. Study 2b

3.2.1. Methods

Thirty new adult volunteers (11 men and 19 women, mean age = 21 years, aged between 19 and 24 years) were paid approximately 20 RMB each to participate in this study. All participants were of Han ethnicity, the majority ethnic group in China. The procedure and design were the same as those in Study 2a, except that the hindering actions and words used in Study 1b were presented. Hence, a 2 (agent group: same vs. different) $\times$ 2 (word meaning: hindering action vs. neutral action) within-subject design was used (see Videos S7 and S8 presenting how the agent from the same and different groups acted).

3.2.2. Results and discussion

3.2.2.1. Within-experiment tests: Fig. 3b shows the mean RTs for the trials with correct responses in all conditions. The RT data were submitted to a repeated-measures 2 (agent group: same vs. different) $\times$ 2 (word meaning: hindering action vs. neutral action) ANOVA after removing the trials with incorrect responses (exclusion rate = 3.2%). No effects were significant (agent group: $F(1, 29) = 1.12, p = .299, \eta_p^2 = 0.04$; word meaning: $F(1, 29) = 0.32, p = .575, \eta_p^2 = 0.01$; interaction effect: $F(1, 29) < 0.01, p = .991, \eta_p^2 < 0.01$). The Bayesian factor measuring the difference in the RTs between the hindering and neutral action words was 0.24 (0.21) when the agent was from the same (a different) social group.

3.2.2.2. Between-experiment tests: To further test whether action valence influences implicit action generalization, a 2 (action valence as a between-subject factor: helpful in Study 2a vs. harmful in Study 2b) $\times$ 2 (agent group as a within-subject factor: same vs. different) $\times$ 2 (word meaning as a within-subject factor: valenced action vs. neutral action) ANOVA was conducted for the RT data, and a three-way interaction was found ($F(1, 58) = 4.00, p = .050, \eta_p^2 = 0.06$). This interaction was due to a two-way interaction between the agent group and word meaning when the helpful action was present in Study 2a but not when the harmful action was present in Study 2b. By analyzing the three-way interaction with agent group as a fixed variable, we found that the interaction between action valence and word meaning was significant when the expected agent belonged to the same group as the two exemplars ($F(1, 58) = 4.01, p = .050, \eta_p^2 = 0.07$) but not when this agent was from a different group ($F(1, 58) = 0.20, p = .656, \eta_p^2 < 0.01$). Similar to Experiment 1, we conducted the same analysis, focusing on the significant two-way interaction effect. In Study 2a, we observed that the facilitation effect of the helpful action ($M = 28$ ms; $SE = 9$) was greater than that of the harmful action in Study 2b when the expected agent belonged to the same group as the exemplars ($M = 5$ ms; $SE = 8$; $t(58) = 2.01, p = .050$, Cohen's $d = 0.52$, 95% CI of the mean difference = [0, 46]).

In summary, increasing the group size from three to five yielded similar results, once again suggesting that helpful actions are more likely to be implicitly generalized among group members than harmful actions.

4. Study 3

Studies 1 and 2 revealed that only helping as a positive helpful action but not hindering as a negative harmful action is implicitly generalized to group members. Although the helping and hindering action words matched in terms of features such as rated familiarity, word frequency, and the number of strokes, they are nonetheless different kinds of words and may have different salience in our minds; even the concepts of the selected hindering action words would not precisely match the expected contents. In this case, even if the hindering action is implicitly expected to be shared among group members, due to specific features of hindering action words, the current paradigm may fail to detect this generalization effect. The current study was designed to test this possibility. Specifically, after two exemplars' actions were observed, the last moving agent was one of the two exemplars (i.e., the old agent) or a new agent, unlike in Study 1. Given that hindering as a negative action can be associated with a particular agent (Alves, Koch, & Unkelbach, 2017a, 2017b; Unkelbach et al., 2020), the same action should be more likely to be expected to be repeated for a particular agent (i.e., within-person action generalization). In this case, if the hindering action words were not implicitly available, we should still observe the null effect for hindering action when expecting the old agent in Study 3; otherwise, contrary to the finding when expecting the new agent in Study 1 (i.e., between-person action generalization), the response to the hindering action words should be facilitated relative to the neutral action words when expecting the old agent. This study also included the condition that the judged agent was a new group member to replicate the findings of Study 1.

4.1. Study 3a

4.1.1. Methods

Thirty new adult volunteers (12 men and 18 women, mean age = 20 years, aged between 18 and 26 years) were paid approximately 20 RMB each to participate in this study. All participants were of Han ethnicity, the majority ethnic group in China. The procedure and design were almost the same as those in Study 1a, except as noted below. In this study, the last agent was always from the same group as the exemplars but was either a new member (i.e., a member who did not demonstrate in the “group members’ action” phase; see Video S9) or the first demonstrating exemplar (i.e., the first agent to perform helping actions in the “group members’ action” phase; see Video S10) with a block design. In this case, the participants completed 104 trials, which followed a 2 (agent source: new vs. old) $\times$ 2 (word meaning: helping action vs. neutral action) within-subject design with 20 trials in each condition and 24 fillers. The order in which blocks presented new or old agents was counterbalanced across participants.

4.1.2. Results and discussion

Fig. 4a shows the mean RTs for the trials with correct responses in all conditions (exclusion rate = 2.1%). Repeated-measures 2 (agent source: new vs. old) $\times$ 2 (word meaning: helping action vs. neutral action) ANOVA was conducted on the RT data. The main effect of agent

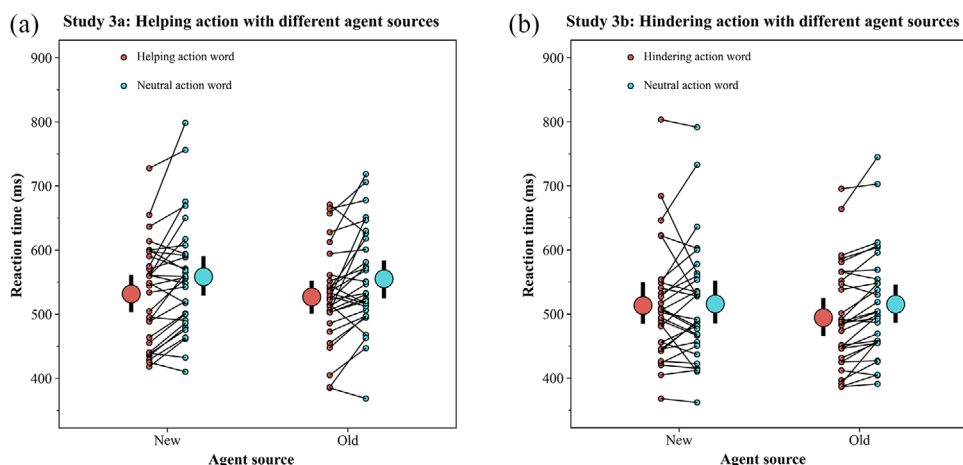


Fig. 4. Mean response time (RT) for the correct responses in (a) Study 3a and (b) Study 3b as a function of the agent source and word meaning. The dots indicate the mean RT for each participant in each condition and are connected for the paired conditions. The large circles denote the means, and the black vertical lines indicate the CIs of the means.

source ($F(1, 29) = 0.12, p = .735, \eta_p^2 < 0.01$) was not significant. However, the main effect of word meaning was ($F(1, 29) = 19.33, p < .001, \eta_p^2 = 0.40$), showing that the participants responded more rapidly to helping ($M = 529$ ms, $SE = 12$) than to neutral action words ($M = 557$ ms, $SE = 14$; 95% CI of the mean difference = $[-40, -15]$). Additionally, this effect was not significantly modulated by the agent source; that is, the interaction effect ($F(1, 29) = 0.04, p = .846, \eta_p^2 < 0.01$) was not significant. The Bayesian factors measuring the difference in the RTs between the helping and neutral action words were 27.99 when the agent was new and 26.53 when the agent had been previously encountered. Hence, this study not only shows that helpful action can be linked with a particular agent, that is, within-person implicit generalization, but also suggests that this type of action can implicitly generalize between group members, which replicates the results of Study 1a.

4.2. Study 3b

4.2.1. Methods

Thirty new adult volunteers (20 men and 10 women, mean age = 21 years, aged between 18 and 27 years) were paid approximately 20 RMB each to participate in this study. All participants were of Han ethnicity, the majority ethnic group in China. The procedure and design were the same as those in Study 3a, except that the hindering actions and words used in Study 1b were presented. Hence, a 2 (agent source: new vs. old) $\times$ 2 (word meaning: hindering action vs. neutral action) within-subject design was obtained (see Videos S11 and S12 presenting how the new and old agents acted).

4.2.2. Results and discussion

4.2.2.1. Within-experiment tests: Fig. 4b shows the mean RTs for trials with correct responses in all conditions (exclusion rate = 2.7%). The repeated-measures 2 (agent source: new vs. old) $\times$ 2 (word meaning: hindering action vs. neutral action) ANOVA was conducted on the RT data. The main effect of the agent source ($F(1, 29) = 1.45, p = .238, \eta_p^2 = 0.05$) was not significant, but the main effect of the word meaning was ($F(1, 29) = 4.40, p = .045, \eta_p^2 = 0.13$), showing that the participants responded more rapidly to the hindering ($M = 504$ ms, $SE = 15$) than to the neutral action words ($M = 516$ ms, $SE = 16$; 95% CI of the mean difference = [1, 23]). Importantly, the interaction effect between the two factors ($F(1, 29) = 5.95, p = .021, \eta_p^2 = 0.17$) was significant. The simple effects analysis revealed that when the last agent had previously been encountered, the participants responded more quickly when the words were labeled a hindering interaction ($M = 494$ ms, $SE = 15$) than when they were labeled a neutral interaction ($M = 515$ ms, $SE = 15$; $t(29) = 5.07, p < .001, d = 0.93$, 95% CI of the mean difference = [-29, -12]). When the last agent was a new group member, there was no significant difference in the RTs between the hindering ($M = 514$ ms, $SE = 17$) and neutral action words ($M = 516$ ms, $SE = 17$; $t(29) = 0.28, p = .782$, Cohen's $d = 0.05$, 95% CI of the mean difference = [-20, 15]). The Bayesian factors measuring the difference in the RTs between the hindering and neutral action words were 0.20 when the agent was new and 1080.70 when the agent had been previously encountered. Hence, the hindering action is implicitly generalized only at the within-person level, suggesting that the harmful action can be associated with a particular individual and that this generalization can be inspected by the current paradigm. At the between-person level, it is not present. These findings ruled out the possibility that the hindering action words were inappropriately selected.

4.2.2.2. Between-experiment tests: A 2 (action valence as a between-subject factor: helpful in Study 3a vs. harmful in Study 3b) $\times$ 2 (agent source as a within-subject factor: new vs. old) $\times$ 2 (word meaning as a within-subject factor: valenced action vs. neutral action) ANOVA was conducted for the RT data, and a three-way interaction did not reach significance ($F(1, 58) = 2.04, p = .159, \eta_p^2 = 0.03$). However, a 2 (action valence as a between-subject factor: helpful in Study 3a vs. harmful in Study 3b) $\times$ 2 (word meaning as a within-subject factor: valenced action vs. neutral action) ANOVA revealed a significant two-way interaction when the agent was new ($F(1, 58) = 4.57, p = .037, \eta_p^2 = 0.07$). The results showed that when the two exemplars performed helpful actions ($M = 27$ ms; $SE = 7$), the facilitation effect was greater than when they performed harmful actions ($M = 2$ ms; $SE = 9$; $t(58) = 2.14, p = .037$, Cohen's $d = 0.55$, 95% CI of the mean difference = [2, 47]); hence, the results replicated those of Study 1. Thus, valence-dependent implicit action generalization was documented again, and implicit generalization does not come as easily for harmful actions as for helpful actions.

However, in the case of an old agent where action generalization occurs at the within-person level, the ANOVA revealed that the two-way interaction was not statistically significant ($F(1, 58) = 0.70, p = .408, \eta_p^2 = 0.01$). This indicates that the varying effects of generalizing helpful and harmful actions at the between-person level (i.e., action generalization among group members) were not a result of unmatched features between helping-related and

hindering-related words. This is because the detected facilitation effect remains nearly identical when the action is generalized at the within-person level. Moreover, as reported in the Supplementary Material, the congruence between the meaning of helping action words and the concept of “help” was matched with the congruence between the meaning of hindering action words and the concept of “hinder.” Hence, the asymmetric generalization was not due to the different saliences or congruences of the words and concepts associated with helping actions from those of the words and concepts associated with hindering actions.

4.2.3. *Replication of findings regarding hindering actions*

Previous findings have indicated that implicit generalization occurs more easily for helpful actions than for harmful actions. This conclusion was based on the nonsignificant facilitation effect of hindering action, especially when the new agent was a member of the same group as the two exemplars. However, this null effect may be specific to the condition when it was mixed with other comparisons. To further examine the replicability of this null effect, we conducted another study including only the critical condition that the last agent was a new group member, as Study 3b set the agent source as new. We recruited 30 new adult volunteers (16 men and 14 women, mean age = 21 years, aged between 18 and 26 years) to participate in this study. We found no significant difference in the RTs of the hindering ($M = 625$ ms, $SE = 8$) and neutral action words ($M = 618$ ms, $SE = 18$; $t(29) = 0.22$, $p = .620$, Cohen’s $d = 0.09$, 95% CI of the mean difference = $[-20, 33]$). The Bayesian factor measuring this difference was 0.22, which suggested substantial evidence against between-condition differences. Hence, it again replicated the null facilitation effect for hindering action and demonstrated that the participants had difficulty implicitly generalizing a harmful action within group members.

5. Study 4

The finding that perceivers do not generalize shared helpful actions to members of another group suggests that implicit generalization for helpful action may depend on social groups. To more closely examine the role of group membership in action generalization, we designed Study 4 to eliminate social groupness by presenting agents as spatially separated and independent in terms of movement and by using the same visual agents as in Study 1a. This manipulation has been shown to sufficiently weaken social group information (Duan et al., 2021; Yin et al., 2022). Therefore, if group membership is a determining factor for implicit action generalization, a smaller facilitation effect or even a null effect would be expected in this study compared to Study 1a; otherwise, we would expect to see a similar effect to that found in Study 1a. By comparing the results of Study 4 with those of Study 1a, we aimed to gain further insights into the role of social groupness in implicit generalization for helpful action.

5.1. Methods

Thirty new adult volunteers (12 men and 18 women, mean age = 20 years, aged between 18 and 26 years) were paid approximately 20 RMB each to participate in this study. All participants were of Han ethnicity, the majority ethnic group in China. The procedure and design were the same as those in Study 1a, except that some social group cues were removed. Specifically, six cartoon agents moved independently rather than synchronously and then spread out into an ABABAB arrangement instead of the AAABBB arrangement used in Study 1a (see Videos S13 and S14). Consequently, the two loose groups could be defined based on their distinct appearance (i.e., three agents had the same identity). This operation was not aimed at making social groups disappear but rather, on a continuous scale, at reducing their intensity.

Previous studies have documented the validity of the above manipulation (Duan et al., 2021; Yin et al., 2022). To further confirm whether the operation was valid in this study, we recruited 40 new adults (16 men and 24 women, mean age = 22 years, aged between 19 and 26 years) to watch two videos from the “introducing social groups” phase of Study 1a and Study 4 and evaluated whether agents with the same identity looked like a unified group using two items. These two items (“To what extent are the three agents with the same appearance closely connected?” and “To what extent are the three agents with the same appearance a unit?”) were established according to the definition and characteristics of a real social group (Morewedge, Chandler, Smith, Schwarz, & Schooler, 2013), referred to previous studies (Duan et al., 2021; Yin et al., 2022) and were measured on a 7-point scale (1 = not at all to 7 = very much). The internal consistency (Pearson’s correlation) of two items was 0.40 ($p = .011$) for the video from Study 1a and 0.60 ($p < .001$) for the video from Study 4, and a paired t test revealed that social groupness was judged more strongly for the video from Study 1a ($M = 5.66$, $SE = 0.20$) than for the video from Study 4 ($M = 4.35$, $SE = 0.16$; $t(39) = 6.76$, $p < .001$, Cohen’s $d = 1.07$, 95% CI of difference = [0.92, 1.71]). These results suggest that the manipulation of social group information was valid.

5.2. Results and discussion

5.2.1. Within-experiment tests

Fig. 5 shows the mean RTs for the trials with correct responses in all conditions (exclusion rate = 1.1%). A repeated-measures 2 (agent group: same vs. different) $\times$ 2 (word meaning: helping action vs. neutral action) ANOVA was conducted on the RT data. No effects were significant (agent group: $F(1, 29) = 0.80$, $p = .379$, $\eta_p^2 = 0.03$; word meaning: $F(1, 29) = 1.20$, $p = .282$, $\eta_p^2 = 0.04$; interaction effect: $F(1, 29) < 0.01$, $p = .948$, $\eta_p^2 < 0.01$). The Bayesian factors measuring the difference in the RTs of the helping and neutral action words were 0.24 (0.38) when the agent was from the same (a different) group.

5.2.2. Between-experiment tests

To further examine the role of social group information with respect to implicit action generalization when a helpful action was present, a 2 (social group information as a between-subject factor: strong in Study 1a vs. weak in Study 4) $\times$ 2 (agent group as a within-subject factor: same vs. different) $\times$ 2 (word meaning as a within-subject factor: valenced action vs.

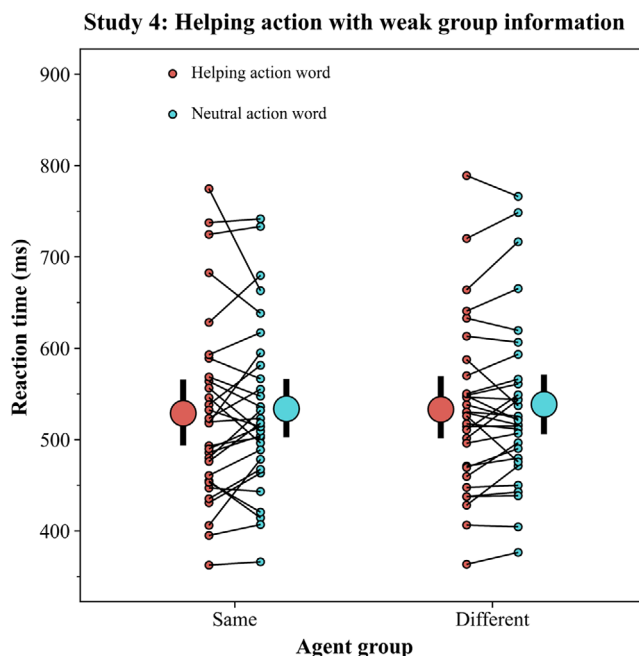


Fig. 5. Mean response time (RT) for correct responses in Study 4 as a function of the agent group and word meaning after removing the group cues. The dots indicate the mean RTs for each participant in each condition and are connected for the paired conditions. The large circles denote the means, and the black vertical lines indicate the CIs of the means.

neutral action) ANOVA was conducted for the RT data, and a three-way interaction was found ($F(1, 58) = 4.24, p = .044, \eta_p^2 = 0.07$). This interaction was due to a two-way interaction between the agent group and word meaning when the social group information was strong in Study 1a but not when it was present in Study 4, as described above. Moreover, we analyzed the three-way interaction by fixing the variable for the agent group and found that the interaction between the social group information and word meaning was significant when the expected agent was from the same group as the two exemplars ($F(1, 58) = 4.95, p = .030, \eta_p^2 = 0.08$) but not when this agent was from a different group ($F(1, 58) < 0.01, p = .971, \eta_p^2 < 0.01$). When the social group information was strong in Study 1a ($M = 31$ ms; $SE = 9$), the facilitation effect was greater for a member of the same group than when the information was weak in Study 4 ($M = 5$ ms; $SE = 8$; $t(58) = 2.22, p = .030$, Cohen's $d = 0.57$, 95% CI of the mean difference = [3, 49]). Hence, group membership (i.e., social category) plays an important role in forming implicit expectations of helpful actions.

6. Study 5

The previous experiments provide evidence that people implicitly tend to generalize helpful actions rather than harmful actions to group members. However, it remains unclear whether

this generalization is solely based on genuine group-based inferences (i.e., sampling actions of different group members), as it might also involve individual-based inferences, where the generalization occurs only when a group member performs the helpful action, regardless of how the action is sampled. However, the sample size is crucial to the inductive generalization of group-based inference (Hendrickson, Perfors, Navarro, & Ransom, 2019; Lawson & Fisher, 2011), as seeing two individuals do the same thing at the same time is a strong indication that this action is shared by group members (Haun, Rekers, & Tomasello, 2012; Herrmann, Legare, Harris, & Whitehouse, 2013). Therefore, the individual-based inference predicted that the implicit generalization of the helping action would not be influenced by how the action is sampled. Group-based inference would imply that the generalization weakens or disappears when the action is repeated by the same group member instead of the majority.

6.1. Methods

Thirty new adult volunteers (12 men and 18 women, mean age = 20 years, aged between 18 and 26 years) were paid approximately 20 RMB each to participate in this study. All participants were of Han ethnicity, the majority ethnic group in China. The stimuli and procedures applied were the same as those used in Study 1a, except as noted below. The condition of the different agent was eliminated, and during the exemplars' demonstration phase, in the condition of the majority demonstration (see Video S15), two agents from the same group performed the helping action. Additionally, in the repetition demonstration condition (see Video S16), the first agent remained still, while the second agent completed the helping action twice. Hence, a 2 (agent demonstration: majority vs. repetition) $\times$ 2 (word meaning: helping action vs. neutral action) within-subject design was used.

6.2. Results and discussion

Fig. 6 shows the mean RTs for trials with correct responses in all conditions (exclusion rate = 1.0%). The RT data were subjected to a repeated-measures 2 (agent demonstration: majority vs. repetition) $\times$ 2 (word meaning: helping action vs. neutral action) ANOVA. The main effect of the agent demonstration was not significant ($F(1, 29) = 0.13, p = .723, \eta_p^2 < 0.01$), but that of the word meaning was significant ($F(1, 29) = 33.74, p < .001, \eta_p^2 = 0.54$), showing that the participants responded more rapidly in the helping word condition ($M = 515$ ms, $SE = 17$) than in the neutral action word condition ($M = 538$ ms, $SE = 18$; 95% CI of the mean difference = [15, 31]). More importantly, the interaction effect was significant ($F(1, 29) = 8.99, p = .006, \eta_p^2 = 0.24$). The simple effects analysis revealed that in the majority demonstration condition, the participants responded more quickly when the meaning of words concerned a helpful interaction ($M = 510$ ms, $SE = 17$) than when their meaning concerned a neutral interaction ($M = 545$ ms, $SE = 19$; $t(29) = 5.40, p < .001$, Cohen's $d = 0.99$, 95% CI of the mean difference = [-47, -21]). The repetition demonstration condition also significantly differed in the RTs between helping ($M = 520$ ms, $SE = 17$) and neutral action words ($M = 532$ ms, $SE = 18$; $t(29) = 2.50, p = .018$, Cohen's $d = 0.46$, 95% CI of the mean difference = [-20, -2]). However, the difference was larger in the majority ($M = 34$ ms, $SE = 6$) than in the repetition demonstration condition ($M = 11$ ms, $SE = 4$; $t(29) =$

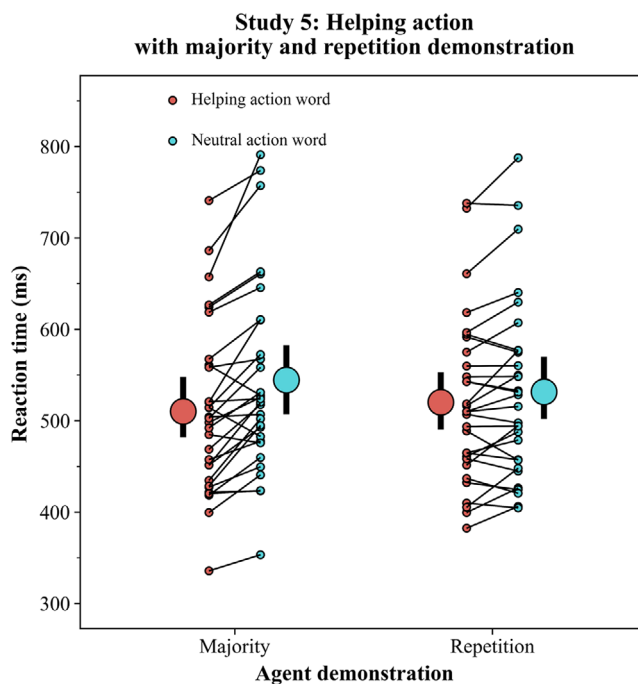


Fig. 6. Mean response time (RT) for correct responses in Study 5 as a function of agent demonstration and word meaning. The dots indicate the mean RTs for each participant in each condition and are connected for the paired conditions. The large circles denote the means, and the black vertical lines indicate the CIs of the means.

3.00, $p = .006$, Cohen's $d = 0.55$, 95% CI of the mean difference = [7, 39]). The Bayesian factors measuring the difference in the RTs for the helping and neutral action words were 2491.89 when the majority performed the helping action and 2.70 when an agent completed the helping action twice. Based on the Bayesian factors, the difference in RTs was not as inclusive in the repetition demonstration condition. In any case, the implicit generalization of helpful action was largely driven by the demonstration of the action completed by multiple group members, thus suggesting that implicit generalization for helpful action relies on group-based inference.

7. Mini meta-analysis

The results of our studies confirm the existence of valence-dependent implicit action generalization. We summarized our findings across all studies (see Table 1). To gain further insight into our findings and to assess the overall size and reliability of our effects, we conducted a mini meta-analysis on the interesting question of whether the RTs for the valenced action words and neutral action words were different (i.e., facilitation effect) when the new agent was a member of the same group as the two exemplars. The meta-analysis was conducted using

Table 1
Overview of studies

Study	Purpose	Design	Results
Study 1 Study 1a ($N = 30$) Study 1b ($N = 30$)	To assess whether the helpful or harmful actions of known group members would be implicitly generalized to an unknown member	Study 1a: Agent group (same vs. different) $\times$ Word meaning (helping action vs. neutral action) Study 1b: Agent group (same vs. different) $\times$ Word meaning (hindering action vs. neutral action)	Helpful actions are implicitly expected to be shared within social groups (Study 1a) but no similar effect was observed for harmful actions (Study 1b)
Study 2 Study 2a ($N = 30$) Study 2b ($N = 30$)	To replicate Study 1's findings with groups of five members	Study 2a: Agent group (same vs. different) $\times$ Word meaning (helping action vs. neutral action) Study 2b: Agent group (same vs. different) $\times$ Word meaning (hindering action vs. neutral action)	The finding of Study 1 was replicated by increasing the group size from three to five
Study 3 Study 3a ($N = 30$) Study 3b ($N = 30$)	To test the possibility that the current paradigm may fail to detect the implicit generalization effect due to specific features of hindering action words	Study 3a: Agent source (new vs. old) $\times$ Word meaning (helping action vs. neutral action) Study 3b: Agent source (new vs. old) $\times$ Word meaning (hindering action vs. neutral action)	Helping actions are implicitly generalized for both new (i.e., a member who did not demonstrate actions) and old (i.e., a member who did demonstrate actions) agents; hindering actions are implicitly generalized only for the old agent
Study 4 ($N = 30$)	To more closely examine the role of group membership in implicit action generalization	Agent group (same vs. different) $\times$ Word meaning (helping action vs. neutral action)	The implicit generalization of helpful actions is specific to the presence of the social group
Study 5 ($N = 30$)	To explore whether the implicit action generalization relies on group- rather than individual-based inference	Agent demonstration (majority vs. repetition) $\times$ Word meaning (helping action vs. neutral action)	The implicit generalization of helpful actions is largely driven by the demonstration of the action completed by multiple group members, rather than being repeated by the same group member

the restricted ML (Maximum likelihood) method in a random effects model in JASP software (JASP Team, 2021; Version 0.16). Concerning helping action, the overall effect across the pilot study, Study 1a, Study 2a, Study 3a, and Study 5 ($N = 140$) was large (Cohen's $d = 0.74$, 95% CI = [0.54, 0.95], $z = 7.12$, $p < .001$). In contrast, for the hindering action across Study 1b, Study 2b, Study 3b, and the replication study ($N = 120$), the overall effect size was trivial (Cohen's $d = 0.07$, 95% CI = [-0.12, 0.26], $z = 0.75$, $p = .457$). When the data on helping action and hindering action were combined, the overall effect was still medium in size (Cohen's $d = 0.43$, 95% CI = [0.18, 0.68], $z = 3.39$, $p < .001$), but the heterogeneity across effect sizes was large ($Q(8) = 24.38$, $p = .002$). Subgroup analysis revealed that the effect sizes varied significantly between helping and hindering actions ($Q(1) = 21.94$, $p < .001$), suggesting that the action valence moderated the effect sizes. The mini meta-analysis further supported the valence-dependent implicit action generalization that helpful actions are more likely than harmful actions to be implicitly generalized within group members.

8. General discussion

The current research explored whether action valence influences implicit action generalization among group members. The results indicate that when known group members performed positive helpful (i.e., helping) actions, words with helping action meanings were identified more quickly than words without such meanings while expecting the possible actions of a new group member from the same group. However, this facilitation effect was not present when the expected new agent was from a different group. Additionally, no such effect was found when known group members performed negative harmful (i.e., hindering) actions, showing that helpful actions are more likely to be implicitly generalized to group members than harmful actions (Study 1). The findings were replicated in Study 2 by expanding the group size. Study 3 further ruled out the possibility that the selected hindering action words did not precisely match the expected contents or had a different salience from that of the helping actions; therefore, no facilitation effect was present. Moreover, the generalization of helpful actions was specific to group membership (Study 4), and the demonstration of helpful actions completed by multiple group members rather than helpful actions repeated by one group member showed group-based inductive inference (Study 5). Hence, it is easier to implicitly generalize helpful than harmful actions within group members, and helpful actions have a privileged status as an attribute shared among group members, supporting valence-dependent implicit action generalization.

8.1. Generalizing actions with social goals

The current findings align with previous studies demonstrating that group members are implicitly expected to act alike but usually with nonsocial goals (Duan et al., 2021; Kalish, 2012; Yin et al., 2022). The present research expands the scope of generalized actions to those with social goals, at least in the case of helping actions. Such expansion suggests that when making inferences about unknown members of social groups, people implicitly generalize

the properties of known members to them and accordingly expect group members to behave consistently and engage in common actions, which is similar to the inductive inference of natural kinds, for example, animals. Nevertheless, for social kinds, inductive generalization may be constrained by socially relevant knowledge (Rakoczy, 2015), as hindering actions that were harmful were found to be more difficult to generalize than helping actions. Moreover, the difficulty in implicitly generalizing harmful actions cannot be explained by the fact that our paradigm was unable to detect the implicit generalization of harmful actions because such information was generalized within a particular agent.

Our findings seem to be inconsistent with studies of stereotyping, which have demonstrated that negative traits can be also generalized across group members. However, the accumulated evidence is insufficient to determine whether negative traits are less generalized among group members than positive traits (Mazziotta, Rohmann, Wright, De Tezanos-Pinto, & Lutterbach, 2015; Paolini, Harwood, & Rubin, 2010; Stark, Flache, & Veenstra, 2013). Moreover, negative stereotypes are usually attributed to an outgroup that may threaten participants who belong to another group (Amodio, 2009; Essien et al., 2017; Miller, Zielaskowski, & Plant, 2012; Sim, Correll, & Sadler, 2013). Thus, group affiliation could motivate participants to hold strong negative prior attitudes toward outgroup members. Nevertheless, the current research focused on the observer's perspective, in which the participants did not belong to any observed group. Hence, the current findings are largely due to the cognitive processing of positive and negative attributes from an observer's perspective but without strong motivations to avoid the targets, at least for helpful and harmful actions. Importantly, the stereotypes are suggested to build upon links between groups and representations of stable traits (Kruse & Degner, 2021; Westra, 2019). While trait generalization does not always imply action generalization and vice versa, action generalization may not occur at the level of trait representation but rather through inferences between groups and specific action-situation combinations.

This asymmetrical generalization of helpful and harmful actions adds to the existing evidence of asymmetric genetic attribution and the updating of representations of such actions (Lebowitz, Tabb, & Appelbaum, 2019; Siegel, Mathys, Rutledge, & Crockett, 2018), overall suggesting that actions with different valences are processed differently. This asymmetrical generalization appears to contradict the evidence suggesting a negativity bias in impression formation for specific individuals. However, such bias emphasizes that negative information is more easily generalized than positive information for individuals across situations rather than across members of social groups (Baumeister et al., 2001; Pratto & John, 1991). In fact, our conclusion aligns with the negativity bias found in impression formation for individuals to some extent. This is because a negative attribute is more likely to be seen as diagnostic of individual characteristics, thereby being particularly associated with that specific individual, as revealed in Study 3b, supporting within-person generalization.

The present conclusion does not preclude the nongeneralization of harmful actions. Rather, it mainly demonstrates that helpful action is more easily implicitly generalized than harmful action, at least in computer-animated scenes. Interestingly, the explicit measurement also indicates that helpful action has a privileged role in action generalization (please see Study S3 in the Supplementary Material). In any case, this study emphasizes implicit and automatic generalization, and how explicit and implicit action generalizations are related and differen-

tiated is beyond the scope of the current research. These topics should be addressed in future research.

8.2. Possible cognitive mechanisms of valence-dependent action generalization

Our findings are consistent with the evaluative information ecology (EvIE) model, which uses external ecology as an explanatory construct and assumes that positive information is generally more frequent and more similar than negative information, thus predicting that relatively positive attributes are judged as being easily shared (Unkelbach et al., 2019). However, this ecological approach should be complemented with category-based inference. Specifically, participants update their beliefs about group attributes similarly for differently valenced actions. However, when inferring a new agent's actions based on formed group attributes, the knowledge of situational norms approved by the group may restrict the inductive generalization of harmful actions more than helpful actions. This is because helpful actions are generally accepted and promoted by social groups, while harmful actions are often prohibited to a greater extent (Ahmed et al., 2020; Gilbert & Basran, 2019). Therefore, people may consider the constraints for the category-based inference of helpful and harmful actions differently, resulting in valence-dependent implicit action generalization.

Admittedly, harmful actions toward perceived threats or “bad” outgroup members may be valued by actors' groups and treated as positive. However, the valence of such actions is often treated as positive, as they are attributed to positive intentions such as protecting group members or punishing violators, despite their negative outcomes. In these situations, harmful actions may also be implicitly generalized, as they fit the boundary knowledge of group-approved situational norms. Hence, further studies could ask participants to learn whether target groups approve of situational norms for harmful actions, and if our speculation is reasonable, knowledge of situational norms should modulate implicit action generalization, especially for harmful actions.

In any case, the current research has revealed an interesting phenomenon: in the implicit process, helpful actions have a privileged status as an attribute in category-based inference. The possible underlying cognitive mechanisms are still under examination.

8.3. Limitations

The current research used computer animation with nonhuman agents. This method allowed possible confounding factors to be strictly controlled for, further allowing us to effectively inspect the rules of generalizing actions. It is crucial to acknowledge the potential limitations of such research in terms of generalizing the findings to real-life situations. The external generalizability of our findings should be examined with recorded human actions in the future. Furthermore, helpful and harmful actions were operationalized as helping and hindering actions, but there were many instantiations of such actions, such as donating and bullying. Additionally, due to space limitations in the visual scene, the current study used only three or five agents for each group, and group size usually plays an important role in group processing (Gerard, Wilhelmy, & Conolley, 1968). With a large sample of hindering actions, the

generalization of such actions may have been present as well. Future studies should use more exemplars and adopt other instantiations of helpful and harmful actions to test our findings.

9. Conclusion

To conclude, helpful actions are easier to implicitly generalize among group members than harmful actions, implying that helpful actions have a privileged status as an attribute shared among group members. This finding shows implicit valence-dependent action generalization, at least in computer-animated scenes. This valence-dependent implicit action generalization highlights that the category-based inference for social action is influenced by action valence, which is associated with distinct knowledge, instead of being solely driven by accumulated evidence of shared attributes across different instances. This is the possible cognitive root explaining why normally promoted actions spread widely and expands our knowledge of social groups.

Author contributions

Jiecheng Huangliang: Conceptualization, formal analysis, investigation, methodology, visualization, writing—original draft, writing—review and editing; Yinfeng Hu, Xutao Zheng, Zikai Xu, and Weiyang Zhou: Conceptualization, formal analysis, investigation, methodology, writing—original draft; Jun Yin: Conceptualization, data curation, funding acquisition, supervision, validation, writing—original draft, writing—review and editing.

Funding

This research was supported by the National Natural Science Foundation of China (Grant No. 32371090).

Note

- 1 We also collected some results of explicit measurements. Given that implicit action generalization is our focus, we reported these results in the Supplementary Material to make the logic in the main context concise.

References

- Ahmed, S., Foulkes, L., Leung, J. T., Griffin, C., Sakhardande, A., Bennett, M., Dunning, D. L., Griffiths, K., Parker, J., Kuyken, W., Williams, J. M. G., Dalgleish, T., & Blakemore, S. J. (2020). Susceptibility to prosocial and antisocial influence in adolescence. *Journal of Adolescence*, 84, 56–68.
- Alves, H., Koch, A., & Unkelbach, C. (2017a). The “common good” phenomenon: Why similarities are positive and differences are negative. *Journal of Experimental Psychology: General*, 146(4), 512–528.

- Alves, H., Koch, A., & Unkelbach, C. (2017b). Why good is more alike than bad: Processing implications. *Trends in Cognitive Sciences*, 21(2), 69–79.
- Alves, H., Koch, A., & Unkelbach, C. (2018). A cognitive-ecological explanation of intergroup biases. *Psychological Science*, 29(7), 1126–1133.
- Amodio, D. M. (2009). Intergroup anxiety effects on the control of racial stereotypes: A psychoneuroendocrine analysis. *Journal of Experimental Social Psychology*, 45(1), 60–67.
- Anderson, S. F., Kelley, K., & Maxwell, S. E. (2017). Sample-size planning for more accurate statistical power: A method adjusting sample effect sizes for publication bias and uncertainty. *Psychological Science*, 28(11), 1547–1562.
- Banaji, M. R., & Bhaskar, R. (2000). Implicit stereotypes and memory: The bounded rationality of social beliefs. In D. L. Schacter & E. Scarry (Eds.), *Memory, brain, and belief* (pp. 139–175). Cambridge, MA: Harvard University Press.
- Baumeister, R. F., Bratslavsky, E., Finkenauer, C., & Vohs, K. D. (2001). Bad is stronger than good. *Review of General Psychology*, 5(4), 323–370.
- Billig, M., & Tajfel, H. (1973). Social categorization and similarity in intergroup behaviour. *European Journal of Social Psychology*, 3(1), 27–52.
- Bodenhausen, G. V., Kang, S. K., & Peery, D. (2012). Social categorization and the perception of social groups. In S. T. Fiske & C. N. Macrae (Eds.), *The SAGE handbook of social cognition* (pp. 311–329). Sage.
- Brainard, D. H. (1997). The Psychophysics Toolbox. *Spatial Vision*, 10(4), 433–436.
- Brewer, M. B., Hong, Y.-Y., & Li, Q. (2004). Dynamic entitativity: Perceiving groups as actors. In V. Yzerbyt, C. M. Judd, & O. Corneille (Eds.), *The psychology of group perception: Perceived variability, entitativity, and essentialism* (pp. 25–38). Psychology Press.
- Bright, A. K., & Feeney, A. (2014). The engine of thought is a hybrid: Roles of associative and structured knowledge in reasoning. *Journal of Experimental Psychology: General*, 143(6), 2082–2102.
- Campbell, D. T. (1958). Common fate, similarity, and other indices of the status of aggregates of persons as social entities. *Systems Research and Behavioral Science*, 3(1), 14–25.
- Castner, J. E., Chenery, H. J., Copland, D. A., Coyne, T. J., Sinclair, F., & Silburn, P. A. (2007). Semantic and affective priming as a function of stimulation of the subthalamic nucleus in Parkinson's disease. *Brain*, 130(5), 1395–1407.
- Chan, W., McCrae, R. R., De Fruyt, F., Jussim, L., Löckenhoff, C. E., De Bolle, M., Costa P. T., Sutin A. R., Realo A., Allik J., Nakazato K., Shimonaka Y., Hřebíčková M., Graf S., Yik M., Brunner-Sciarra M., De Figueiroa N. L., Schmidt V., Ahn C. K., & Terracciano, A. (2012). Stereotypes of age differences in personality traits: Universal and accurate? *Journal of Personality and Social Psychology*, 103(6), 1050–1066.
- Chen, J. M., & Ratliff, K. A. (2015). Implicit attitude generalization from Black to Black–White biracial group members. *Social Psychological and Personality Science*, 6(5), 544–550.
- Crawford, M. T., Sherman, S. J., & Hamilton, D. L. (2002). Perceived entitativity, stereotype formation, and the interchangeability of group members. *Journal of Personality and Social Psychology*, 83(5), 1076–1094.
- Dienes, Z. (2014). Using Bayes to get the most out of non-significant results. *Frontiers in Psychology*, 5, 781.
- Duan, J., Jiang, Y., He, Y., Zhang, F., Shen, M., & Yin, J. (2021). Action generalization across group members: Action efficiency matters. *Cognitive Science*, 45(4), 1–27.
- Essien, I., Stelter, M., Kalbe, F., Koehler, A., Mangels, J., & Meli, S. (2017). The shooter bias: Replicating the classic effect and introducing a novel paradigm. *Journal of Experimental Social Psychology*, 70, 41–47.
- Fazio, R. H., Pietri, E. S., Rocklage, M. D., & Shook, N. J. (2015). Positive versus negative valence: Asymmetries in attitude formation and generalization as fundamental individual differences. In J. M. Olson & M. P. Zanna (Eds.), *Advances in experimental social psychology* (Vol. 51, pp. 97–146). Burlington, VT: Academic Press.
- Frith, C. D., & Frith, U. (2008). Implicit and explicit processes in social cognition. *Neuron*, 60(3), 503–510.
- Gerard, H. B., Wilhelmy, R. A., & Conolley, E. S. (1968). Conformity and group size. *Journal of Personality and Social Psychology*, 8(1p1), 79–82.
- Gilbert, P., & Basran, J. (2019). The evolution of prosocial and antisocial competitive behavior and the emergence of prosocial and antisocial leadership styles. *Frontiers in Psychology*, 10, 610.

- Greenwald, A. G., & Banaji, M. R. (1995). Implicit social cognition: Attitudes, self-esteem, and stereotypes. *Psychological Review*, 102(1), 4–27.
- Greenwald, A. G., Poehlman, T. A., Uhlmann, E. L., & Banaji, M. R. (2009). Understanding and using the Implicit Association Test: III. Meta-analysis of predictive validity. *Journal of Personality and Social Psychology*, 97(1), 17–41.
- Haidt, J., & Kesebir, S. (2010). Morality. In S. T. Fiske, D. Gilbert, & G. Lindzey (Eds.), *Handbook of social psychology* (5th ed., pp. 797–832). Hoboken, NJ: Wiley.
- Hamlin, J. K., Wynn, K., & Bloom, P. (2007). Social evaluation by preverbal infants. *Nature*, 450(7169), 557–559.
- Haun, D. B. M., Rekers, Y., & Tomasello, M. (2012). Majority-biased transmission in chimpanzees and human children, but not orangutans. *Current Biology*, 22(8), 727–731.
- Heit, E., & Rubinstein, J. (1994). Similarity and property effects in inductive reasoning. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 20(2), 411–422.
- Hendrickson, A. T., Perfors, A., Navarro, D. J., & Ransom, K. (2019). Sample size, number of categories and sampling assumptions: Exploring some differences between categorization and generalization. *Cognitive Psychology*, 111, 80–102.
- Herrmann, P. A., Legare, C. H., Harris, P. L., & Whitehouse, H. (2013). Stick to the script: The effect of witnessing multiple actors on children's imitation. *Cognition*, 129(3), 536–543.
- Hill, A. B., & Kemp-Wheeler, S. M. (1989). The influence of context on lexical decision times for emotionally aversive words. *Current Psychology*, 8(3), 219–227.
- Hirschfeld, L. A. (2001). On a folk theory of society: Children, evolution, and mental representations of social groups. *Personality and Social Psychology Review*, 5(2), 107–117.
- Hu, Y., Yang, Y., Huang, P., Ai, D., Sun, H., Zhou, D., Huangliang, J., & Yin, J. (2023). More evidence, greater generalization? The relation between the prevalence of observed action and the strength of generalization depends on action properties. *Journal of Experimental Psychology: Human Perception and Performance*, 49(3), 306–326.
- Isik, L., Koldewyn, K., Beeler, D., & Kanwisher, N. (2017). Perceiving social interactions in the posterior superior temporal sulcus. *Proceedings of the National Academy of Sciences of the United States of America*, 114(43), E9145–E9152.
- Jara-Ettinger, J., Gweon, H., Schulz, L. E., & Tenenbaum, J. B. (2016). The naïve utility calculus: Computational principles underlying commonsense psychology. *Trends in Cognitive Sciences*, 20(8), 589–604.
- JASP Team. (2021). JASP (0.16) [Computer software].
- Jussim, L., Crawford, J. T., & Rubinstein, R. S. (2015). Stereotype (in) accuracy in perceptions of groups and individuals. *Current Directions in Psychological Science*, 24(6), 490–497.
- Kalish, C. W. (2012). Generalizing norms and preferences within social categories and individuals. *Developmental Psychology*, 48(4), 1133–1143.
- Kazak, A. E. (2018). Editorial: Journal article reporting standards. *American Psychologist*, 73(1), 1–2.
- Kemp, C., & Tenenbaum, J. B. (2009). Structured statistical models of inductive reasoning. *Psychological Review*, 116(1), 20–58.
- Kruse, F., & Degner, J. (2021). Spontaneous state inferences. *Journal of Personality and Social Psychology*, 121(4), 774–791.
- Kurdi, B., Mann, T. C., Charlesworth, T. E., & Banaji, M. R. (2019). The relationship between implicit intergroup attitudes and beliefs. *Proceedings of the National Academy of Sciences*, 116(13), 5862–5871.
- Lawson, C. A., & Fisher, A. V. (2011). It's in the sample: The effects of sample size and sample diversity on the breadth of inductive generalization. *Journal of Experimental Child Psychology*, 110(4), 499–519.
- Lebowitz, M. S., Tabb, K., & Appelbaum, P. S. (2019). Asymmetrical genetic attributions for prosocial versus antisocial behaviour. *Nature Human Behaviour*, 3(9), 940–949.
- Liberman, Z., Woodward, A. L., & Kinzler, K. D. (2017). The origins of social categorization. *Trends in Cognitive Sciences*, 21(7), 556–568.
- Lorenz, A. R., & Newman, J. P. (2002). Deficient response modulation and emotion processing in low-anxious Caucasian psychopathic offenders: Results from a lexical decision task. *Emotion*, 2(2), 91–104.

- Ludwig, N. N., Hecht, E. E., King, T. Z., Revill, K. P., Moore, M., Fink, S. E., & Robins, D. L. (2020). A novel social attribution paradigm: The Dynamic Interacting Shape Clips (DISC). *Brain and Cognition*, 138, 105507.
- Macrae, C. N., & Bodenhausen, G. V. (2000). Social cognition: Thinking categorically about others. *Annual Review of Psychology*, 51(1), 93–120.
- Mazziotta, A., Rohmann, A., Wright, S. C., De Tezanos-Pinto, P., & Lutterbach, S. (2015). (How) does positive and negative extended cross-group contact predict direct cross-group contact and intergroup attitudes? *European Journal of Social Psychology*, 45(5), 653–667.
- Miller, S. L., Zielaskowski, K., & Plant, E. A. (2012). The basis of shooter biases: Beyond cultural stereotypes. *Personality and Social Psychology Bulletin*, 38(10), 1358–1366.
- Morewedge, C. K., Chandler, J. J., Smith, R., Schwarz, N., & Schooler, J. (2013). Lost in the crowd: Entitative group membership reduces mind attribution. *Consciousness and Cognition*, 22(4), 1195–1205.
- Navon, M., Shechter, A., & Bar-Anan, Y. (2021). The effect of group membership and individuating information on automatic and deliberate evaluation of well-known people. *Journal of Personality and Social Psychology*, 121(3), 498–523.
- Noyes, A., & Keil, F. C. (2018). Asymmetric mixtures: Common conceptual priorities for social and chemical kinds. *Psychological Science*, 29(7), 1094–1103.
- Paolini, S., Harwood, J., & Rubin, M. (2010). Negative intergroup contact makes group memberships salient: Explaining why intergroup conflict endures. *Personality and Social Psychology Bulletin*, 36(12), 1723–1738.
- Powell, L. J., & Spelke, E. S. (2013). Preverbal infants expect members of social groups to act alike. *Proceedings of the National Academy of Sciences of the United States of America*, 110(41), E3965–E3972.
- Pratto, F., & John, O. P. (1991). Automatic vigilance: The attention-grabbing power of negative social information. *Journal of Personality and Social Psychology*, 61(3), 380–391.
- Rakoczy, H. (2015). Comparative metaphysics: The development of representing natural and normative regularities in human and non-human primates. *Phenomenology and the Cognitive Sciences*, 14(4), 683–697.
- Ranganath, K. A., & Nosek, B. A. (2008). Implicit attitude generalization occurs immediately; Explicit attitude generalization takes time. *Psychological Science*, 19(3), 249–254.
- Ratliff, K. A., & Nosek, B. A. (2011). Negativity and outgroup biases in attitude formation and transfer. *Personality and Social Psychology Bulletin*, 37(12), 1692–1703.
- Robinson-Riegler, G. L., & Winton, W. M. (1996). The role of conscious recollection in recognition of affective material: Evidence for positive-negative asymmetry. *Journal of General Psychology*, 123(2), 93–104.
- Schein, C., & Gray, K. (2018). The theory of dyadic morality: Reinventing moral judgment by redefining harm. *Personality and Social Psychology Review*, 22(1), 32–70.
- Siegel, J. Z., Mathys, C., Rutledge, R. B., & Crockett, M. J. (2018). Beliefs about bad people are volatile. *Nature Human Behaviour*, 2(10), 750–756.
- Sim, J. J., Correll, J., & Sadler, M. S. (2013). Understanding police and expert performance: When training attenuates (vs. exacerbates) stereotypic bias in the decision to shoot. *Personality and Social Psychology Bulletin*, 39(3), 291–304.
- Stark, T. H., Flache, A., & Veenstra, R. (2013). Generalization of positive and negative attitudes toward individuals to outgroup attitudes. *Personality and Social Psychology Bulletin*, 39(5), 608–622.
- Taylor, S. E. (1991). Asymmetrical effects of positive and negative events: The mobilization-minimization hypothesis. *Psychological Bulletin*, 110(1), 67–85.
- Unkelbach, C., Alves, H., & Koch, A. (2020). Negativity bias, positivity bias, and valence asymmetries: Explaining the differential processing of positive and negative information. In B. Gawronski (Ed.), *Advances in experimental social psychology* (pp. 115–187). Cambridge, MA: Academic Press.
- Unkelbach, C., Koch, A., & Alves, H. (2019). The evaluative information ecology: On the frequency and diversity of “good” and “bad”. *European Review of Social Psychology*, 30(1), 216–270.
- Vallat, R. (2018). Pingouin: Statistics in Python. *Journal of Open Source Software*, 3(31), 1026.
- Van Rossum, G., & Drake, F. L. (2009). *Python 3 Reference Manual*. Scotts Valley, CA: CreateSpace.
- Westra, E. (2019). Stereotypes, theory of mind, and the action–prediction hierarchy. *Synthese*, 196(7), 2821–2846.

- Xu, H., Duan, J., Chen, M., & Yin, J. (2019). The mechanisms of action prediction based on social group information. *Chinese Journal of Applied Psychology*, 25(3), 239–252.
- Yin, J., Duan, J., Huangliang, J., Hu, Y., & Zhang, F. (2022). Members of highly entitative groups are implicitly expected to behave consistently based on their deep-level goals instead of their shallow-level movements. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 48(1), 13–28.
- Yin, J., Zhou, D., Ai, D., Sun, H., Duan, J., Sun, Z., & Guo, X. (2023). Event-related potential and behavioural evidence of goal-based expectations for consistent actions among group members. *British Journal of Psychology*, 114(3), 662–677.

Supporting Information

Additional supporting information may be found online in the Supporting Information section at the end of the article.

Supporting Information