

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Effects of Egocentric and Exocentric Strategies on Reciprocity: Agent-Based Study Using Ultimatum Game

Permalink

<https://escholarship.org/uc/item/3fn0z7zb>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 48(0)

Authors

Oasa, Kishin

Shimojo, Shigen

Hayashi, Yugo

Publication Date

2026

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

Effects of Egocentric and Exocentric Strategies on Reciprocity: Agent-Based Study Using Ultimatum Game

Kishin Oasa (cp0150sr@ed.ritsumei.ac.jp)¹, Shigen Shimojo² & Yugo Hayashi³

¹Graduate School of Human Science, Ritsumeikan University

²Research Organization of Open Innovation & Collaboration, Ritsumeikan University

³College of Comprehensive Psychology, Ritsumeikan University

Abstract

This study investigates the formation of reciprocal altruism in human-agent interactions using the ultimatum game, focusing on how people infer others' action-generation principles. We examine the effects of different agent strategies on human rejection behavior and agent impression. An experiment is conducted where two factors are manipulated in the agent's strategy: the presence of a selfish, egocentric strategy based on the agent's internal policy; and an altruistic, exocentric strategy that adjusts to the participant's feedback. The results show that participants rejected proposals less frequently when the agent adopted exocentric/non-egocentric strategies. Furthermore, the exocentric strategy increased the perception of safety toward the agent, whereas the egocentric strategy decreased the perception of intelligence. These findings suggest that cooperative behavior and social impressions are shaped not only by outcome distributions but also by inferences about latent policy parameters, such as responsiveness to social signals and rigidity of internal dynamics.

Keywords: HAI; agent; ultimatum game; reciprocal altruism; other-model formation; intention inference; social contingency

Introduction

In social interactions, humans observe the behavior of others and adjust their own behavior while inferring the underlying internal states, such as beliefs, intentions, policies, or personality. This ability to infer the internal states of others is known as the theory of mind (Premack & Woodruff, 1978). It has been considered to contribute significantly to the development of sociality and as a basis for all social interactions, such as cooperation, competition, and trust formation (Yoshida et al., 2009). In particular, people tend to form models of others based on their behavioral patterns, such as “what policies are they pursuing?” and “How will they react to my actions?”, and then use these models to make adaptive decisions.

Altruism is one of the most important factors in the formation of cooperative relationships. In evolutionary biology, altruistic behavior is defined as an action that incurs a cost to the actor while benefiting another individual. Although such behavior appears maladaptive from the perspective of individual fitness, altruism is widely observed in humans and other species. The theory of kin selection (Hamilton, 1964) explains altruism toward genetic relatives, whereas the theory of reciprocal altruism accounts for altruistic behavior among unrelated individuals by emphasizing the long-term benefits of mutual cooperation (Trivers, 1971).

For reciprocal cooperation to remain stable, individuals must distinguish cooperative partners from selfish free riders (Stephens, 1996). Accordingly, various psychological mechanisms for detecting uncooperative individuals have been proposed (e. g. Fehr and Fischbacher (2003) and Rabin (1993)).

The ultimatum game provides a powerful framework for investigating this type of social interaction. In this classic economic game, a proposer offers a division of resources and a responder decides whether to accept it; if the responder rejects the offer, both parties receive nothing. Although rationality predicts that responders should accept any non-zero offer, empirical studies consistently show that unfair proposals are frequently rejected (Güth et al., 1982). Such rejections are often explained by inequity aversion, i.e., a preference to avoid disadvantageous or unfair outcomes even at a personal cost (Fehr & Schmidt, 1999). Alternatively, they are also interpreted as a form of altruistic punishment, whereby individuals incur costs to enforce social norms of fairness (Fehr & Gächter, 2002).

Crucially, accumulating evidence suggests that partner evaluation is based not only on distributional outcomes but also on the inference of the underlying intentions and behavioral policies of others. For instance, excessively generous behavior in a dictator game—an ultimatum game where the recipient cannot reject the proposal—can elicit suspicion rather than trust, obscuring the actor's intentions (Mochizuki, 2014). Likewise, models of reciprocity propose that humans evaluate not only the fairness of outcomes but also the intentions inferred from others' actions when selecting cooperative partners (Falk & Fischbacher, 2006).

Some previous studies have extended the ultimatum game paradigm to human-agent interaction, examining how humans interpret and respond to artificial agents in social decision-making contexts. For example, Hayashi and Okada (2017) reported that human-like “irrational” decisions emerged when participants believed their opponents were human and when the opponent's behavior was ambiguous. Moreover, Hamada et al. (2023) reported that when agents adopt adaptive strategies, participants perceive greater closeness and show stronger preferences for fair and reciprocal decisions.

Taken together, these findings suggest that humans are sensitive not only to behavioral outcomes but also to the generative processes underlying others' actions. In other words, altruism is not simply a matter of how much one gives, but also of how and why one chooses to give.

This raises a critical question: how do humans infer such intentions and cooperative dispositions from observed behavior? Developmental psychology highlights the importance of social contingency—the degree to which an agent’s behavior is responsive to another’s actions—as a key cue for identifying social agents (György & Watson, 1999; Johnson, 2000; Watson, 1972). Similarly, research on human–robot interaction indicates that robots who predict their human partner’s actions and adapt their strategies accordingly are perceived as more intelligent and trustworthy (Hoffman & Breazeal, 2007). This implies that for intention inference and cooperation, the manner and extent of a partner’s responsiveness is a significant variable.

However, in most economic game studies, altruism has been operationalized almost exclusively in terms of distributional outcomes, such as the amount of resources allocated to the partner. Such static operationalizations overlook the dynamic nature of social interaction and provide limited insight into how humans infer others’ action-generation principles. Consequently, it remains unclear how information about others’ strategic responsiveness is integrated into reciprocal altruistic behavior and social impression formation.

Purpose and Hypotheses

The present study addresses this gap by redefining altruism as responsiveness to a partner’s social signals rather than the amount of resources allocated. Using a repeated ultimatum game, we manipulated whether an agent adjusted its proposals based on the responder’s previous reactions (exocentric strategy) and whether it followed a fixed self-centered policy independent of feedback (egocentric strategy), thereby examining how humans infer and evaluate these latent properties.

Interactions between humans and agents do not necessarily exhibit the same characteristics as those between humans. However, because we tend to unconsciously assign human characteristics to nonhuman entities and regard them as social beings (Nass et al., 1996), these characteristics can serve as a good analogy for understanding the interactions between humans. More importantly, by using the framework of human–agent interaction, one may obtain high-quality data by observing actual human cognition and behavior while manipulating the behavior of others as controllable variables, which can then result in a more precise understanding of the interaction process.

As shown in previous studies, rejection behavior can be understood not only as an outcome of resource distribution but also as a socially meaningful response to perceived intentions and strategies of the proposer. Thus, in the present study, we use rejection rates as an index of reciprocal and altruistic behavior and subjective evaluations of agents as indices of inferred social properties. We hypothesize that agents exhibiting high social responsiveness will be perceived as cooperative partners and elicit fewer rejections, whereas agents exhibiting rigid, self-centered policies will be perceived as less cooperative and elicit more rejections. The questionnaire used

was the Godspeed Questionnaire Series (GQS) developed by Bartneck (2023), which measures five dimensions: anthropomorphism, animacy, likability, perceived intelligence, and perceived safety.

For the experiment, we formulated the following two hypotheses:

H1: Participants will reject fewer proposals from agents adopting an exocentric strategy and more proposals from agents adopting an egocentric strategy.

Furthermore, if humans form multidimensional models of others based on their action-generation processes, strategic responsiveness should selectively affect social impressions beyond mere likability. Adaptive agents are perceived as safer and more reliable interaction partners.

H2: Agents adopting an exocentric strategy and not adopting an egocentric strategy will be rated higher in terms of perceived safety and likability on the GQS.

Method

Participants

A total of 160 participants (51 males, 107 females, and 2 N/A) participated in the experiment. They were randomly assigned to one of four conditions: with or without an exocentric strategy, and with or without an egocentric strategy, each with 40 participants. The average age of the participants was 18.86 years ($SD = 0.77$). This experiment was approved by the ethics committee of the authors’ institution. Although the participants’ age range was narrow, behavioral patterns in the ultimatum game are generally known to remain relatively stable from late adolescence through adulthood (Sutter, 2007); therefore, the age-related constraint in the present sample is unlikely to substantially undermine the validity of this study.

Procedure

The participants wore a head-mounted display (HMD) and played an ultimatum game with an agent in a virtual reality (VR) space using a system developed previously (Kusatake, 2020).

Before the game commenced, the participants were familiarized with the movements in the VR space and the operation of the controller while wearing the HMD. They were provided with an explanation about the game’s content and were informed that the agent in front of them was their opponent. Additionally, they were informed that the experimental data would remain anonymous and that they were allowed to act freely based on their own judgment. The game was repeated 15 times, with the agent and participant serving as the proposer and responder, respectively. After the trial, the participants were instructed to complete a questionnaire survey using the GQS. We used Oculus Rift S (Oculus VR, LLC) as the HMD. For the agent avatar, we used a humanoid model published in the character model-sharing service VRoid Hub (Pixiv Inc.) (Kasaneko, 2019). An example of a participant’s field-of-view during the game is shown in Figure 1.

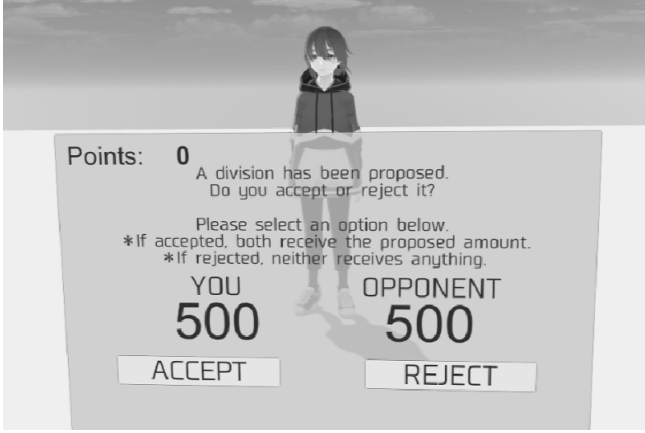


Figure 1: Example of participant view in the experiment.

Table 1: Possible proposals from agent and their share for each party.

proposal options	share	
	participant	agent
0	900	100
1	800	200
2	700	300
3	500	500
4	300	700
5	200	800
6	100	900

Experimental System

The game involved distributing a 1000-point reward between both players. In each round, the agent selected one proposal from seven possible allocation options, each indexed from 0 to 6 as shown in Table 1. In the first round, each agent presented Proposal 3, which corresponded to an equal distribution of 500:500 points. In subsequent rounds, the proposal selection was subject to a constraint: agents can only propose the same option they had selected in the previous round or an option with an adjacent index. Thus, in the second round, the agents selected one of Proposal 2, 3, or 4.

Selecting a proposal that increased the agent’s own share by one level relative to the previous round was defined as the “raise” action (e.g., selecting Proposal 5 after Proposal 4). Selecting the same proposal as in the previous round was defined as “hold,” and selecting a proposal that reduced the agent’s own share by one level was defined as “lower.” Accordingly, the set of possible proposal actions for the agent is expressed as follows:

$$\mathcal{A} = \{\text{raise}, \text{hold}, \text{lower}\} \quad (1)$$

In each round, one of the proposal actions was selected based on the strategy specified for the corresponding experimental condition, and the proposal with the corresponding

index was presented to the participant.

At the beginning of each round, the message “Waiting for a proposal from the computer. Please wait a moment” was displayed for 5, 10, or 15 s. After this waiting period, the allocation proposed by the agent and the buttons “Accept” and “Reject” were displayed and selected by the participants. Throughout the game, the interface continuously displayed the cumulative scores earned by the participants. While waiting, the agent performed a waiting motion by gently swaying its body from side to side.

Agents with an exocentric strategy were designed to increase/decrease behaviors accepted/rejected by the participants. By contrast, agents with an egocentric strategy selected proposal actions based on their own internal policy, independent of the participant’s responses. In the present experiment, these strategies were operationalized by manipulating the information sources referred by the agent when selecting the next proposal action: (a) the participant’s response and (b) the agent’s previous action. By combining the presence or absence of these two information sources, four types of agent strategies were constructed, as described below. Hereafter, we abbreviate the exocentric and egocentric strategies as *exo* and *ego*, respectively. We denote the presence (absence) of each strategy with + (−) (e.g., *exo+ego−*).

exo− / ego− In the non-exocentric/non-egocentric condition, the agent selects its proposal action randomly without depending on either its previous proposal or the participant’s response. Specifically, three possible actions—raise, hold, and lower—are selected with equal probability (33%) in each round.

exo+ / ego− In the exocentric/non-egocentric condition, the agent does not have an internal probabilistic policy and thus selects actions based only on the participant’s response. The agent adopts a win–stay–lose shift strategy (Nowak & Sigmund, 1992): if an action (e.g., raise) is accepted in the previous round, then the agent repeats the same action in the next round; if it is rejected, then the agent switches to the opposite action (e.g., lower). If the previous action is held and the proposal is rejected, then the agent randomly selects to either raise or lower in the next round.

exo− / ego+ In the non-exocentric/egocentric condition, the agent selects its action based solely on its previous action regardless of the participant’s response. Let the proposal action in round t be denoted as $a_t \in \mathcal{A}$. The probability of selecting an action is denoted as $P(a_t|a_{t-1})$ and can be expressed as a transition probability matrix.

$$P(a_t|a_{t-1}) = \begin{pmatrix} 0.116 & 0.130 & 0.754 \\ 0.210 & 0.676 & 0.114 \\ 0.430 & 0.349 & 0.221 \end{pmatrix} \quad (2)$$

For example, when the action in the previous round is raise, then the probability of selecting lower in the next round is

75.4%. This transition matrix was derived from a preliminary experiment in which human participants played the same game with other human players and the proposal behaviors of the proposers were recorded. In the first round, the agent's previous action was defined as hold.

exo+ / ego+ In the exocentric/egocentric condition, the agent determines its action by referring to both its previous action and the participant's response. Initially, the agent follows the same transition probabilities as those in the non-exocentric/egocentric condition. However, as the game progresses, the probability table is dynamically updated through learning, where participant acceptance and rejection are used as a positive and negative reward, respectively. Learning is implemented using the standard Q-learning. A Q-table $Q(s_t, a_t)$ is defined such that the state $s_t \in \mathcal{A}$ corresponded to the agent's previous action ($s_t = a_{t-1}$). The reward r_{t+1} is set to +1 for acceptance and -1 for rejection, and the Q-values are updated based on the following equation, with a discount factor $\gamma = 0.9$ and a learning rate $\alpha = 0.5$:

$$Q(s_t, a_t) \leftarrow Q(s_t, a_t) + \alpha \left[r_{t+1} + \gamma \max_{a' \in \mathcal{A}} Q(s_{t+1}, a') - Q(s_t, a_t) \right] \quad (3)$$

Subsequently, the Q values are converted into action-selection probabilities using the softmax function as follows:

$$p(a_t | s_t) = \frac{\exp Q(s_t, a_t)}{\sum_{a' \in \mathcal{A}} \exp Q(s_t, a')} \quad (4)$$

The Q-table is initialized using the logarithm of the transition probabilities $\ln P(a_t | a_{t-1})$ shown in Equation 2 such that the initial action-selection probabilities match those of the non-exocentric/egocentric strategy. Technically, the algorithms used in the non-exocentric/egocentric strategy are equivalent to those in the exocentric/egocentric strategy, with the learning rate α set to 0.0.

Results

A two-way analysis of variance was conducted for each dependent variable, with the presence or absence of exocentric and egocentric strategies as between-participant factors. The dependent variables were (a) the total number of times the participants rejected the agent's proposals across the 15 rounds, and (b) the mean scores of each subscale of the GQS. The results of statistical tests for differences are shown in the graphs as $p < .05^*$, $p < .01^{**}$, and $p < .10^\dagger$. Error bars in graphs indicate standard error.

Rejections by Participants

Figure 2 presents the mean number of proposal rejections under each experimental condition. The interaction between the exocentric and egocentric strategies was not significant ($F(1, 156) = 0.92, p = .339, \eta_p^2 = 0.006$). However, significant main effects were observed for both fac-

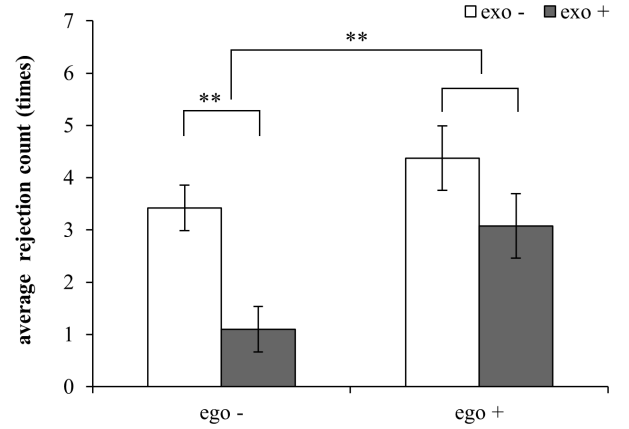


Figure 2: Average number of times participants rejected agent's proposals.

tors. The participants rejected the agent's proposals less frequently when the agent did not adopt an egocentric strategy ($F(1, 156) = 7.50, p = .007, \eta_p^2 = 0.046$) and when the agent adopted an exocentric strategy ($F(1, 156) = 11.52, p = .001, \eta_p^2 = 0.069$). These results indicate that the absence of an egocentric strategy and the presence of an exocentric strategy independently reduced the participants' rejection behavior.

Godspeed Questionnaire

Figures 3–5 present the mean scores for selected subscales of the GQS across conditions.

Likability For the likability subscale (Figure 3), the interaction between the exocentric and egocentric strategies was not significant ($F(1, 156) = 1.86, p = .174, \eta_p^2 = 0.012$). Neither main effects were statistically significant. However, a pairwise comparison using the Holm–Bonferroni method revealed that, when an egocentric strategy was present, agents adopting an exocentric strategy were rated as more likable than those without an exocentric strategy ($t(156) = -2.11, p = .037$).

Perceived Safety In terms of perceived safety (Figure 4), the interaction effect was not significant ($F(1, 156) = 0.06, p = .805, \eta_p^2 = 0.000$). A significant main effect of the exocentric strategy was observed ($F(1, 156) = 7.03, p = .009, \eta_p^2 = 0.043$), thus indicating that the participants perceived agents employing an exocentric strategy as safer.

Perceived Intelligence In terms of perceived intelligence (Figure 5), the interaction effect was not significant ($F(1, 156) = 2.11, p = .149, \eta_p^2 = 0.013$). The egocentric strategy indicated a significant main effect ($F(1, 156) = 4.45, p = .037, \eta_p^2 = 0.028$), where agents adopting an egocentric strategy were rated as less intelligent. Additionally,

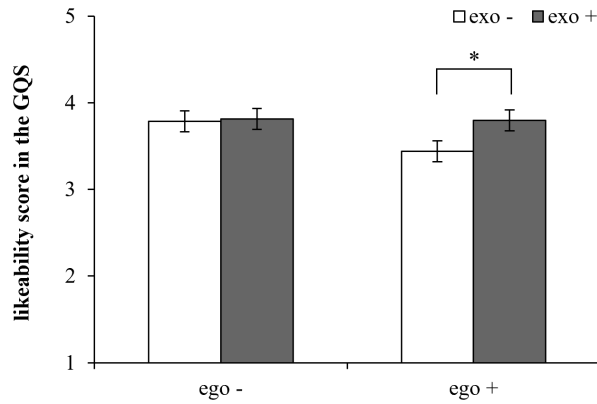


Figure 3: Average score of “likeability” items in the GQS.

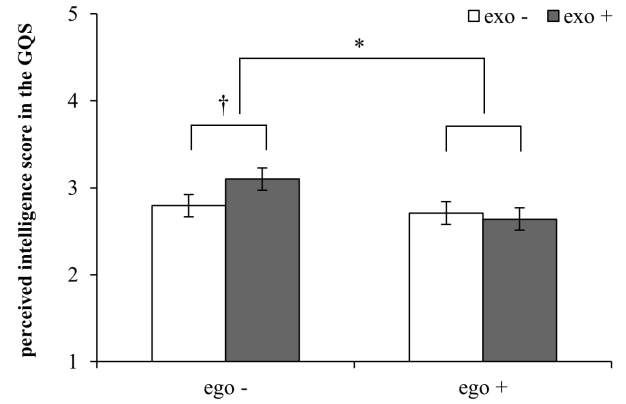


Figure 5: Average score of “perceived intelligence” items in the GQS.

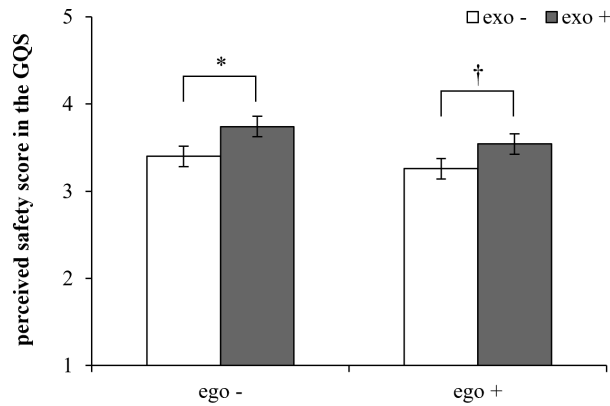


Figure 4: Average score of “perceived safety” items in the GQS.

when the egocentric strategy was absent, agents with an exocentric strategy tended to receive higher intelligence ratings than those without, although this difference was not statistically significant ($t(156) = -1.67, p = .097$).

Discussion

The number of rejected proposals decreased both when the agent employed an exocentric strategy and when it did not employ an egocentric strategy, supporting Hypothesis H1. These results indicate that an agent’s exocentric behavioral policy tends to promote altruistic and reciprocal responses from human participants. When another party voluntarily increases behaviors that are beneficial to the individual, people may infer that the partner possesses a cooperative or altruistic orientation that encourages more cooperative responses.

By contrast, the adoption of an egocentric strategy increased rejection rates. When the agent adhered rigidly to its internal policy and showed limited sensitivity to the participant’s responses, such “stubbornness” may be interpreted not merely

as selfishness but as low social contingency. Under such conditions, participants may have engaged in rejection as a form of norm enforcement or altruistic punishment, aimed at preventing exploitation and restoring fairness in the interaction (Fehr & Gächter, 2002). These results align with theories of reciprocal altruism, which posit that humans selectively cooperate with partners who exhibit signals of future reciprocation.

A trial-level analysis of rejection rates provides further evidence for the dynamic process of intention inference. In both the exo-/ego- and exo-/ego+ conditions, the mean rejection rate rose to approximately 30% during the first two or three rounds. Thereafter, however, the exo-/ego+ condition maintained roughly the same level of rejection, whereas the exo-/ego- condition settled to around 20%. Importantly, this occurred even though, given the settings of the transition probabilities, the exo-/ego+ condition was in fact more likely to generate fairer offers over time. This temporal shift suggests that participants did not merely evaluate static offer distributions; instead, their decision-making was likely influenced by the transition patterns of the proposals.

The GQS results partially supported Hypothesis H2. In terms of perceived safety, a significant main effect of the exocentric strategy was observed: agents that considered participants’ responses were evaluated as safer and more trustworthy. This suggests that responsiveness to social signals reduces perceived risk in future interactions. From the participants’ perspective, an agent that adapts to feedback is likely to be viewed as a predictable and cooperative partner, thereby lowering the perceived necessity of defensive or punitive behavior.

Likability showed a more nuanced pattern. Although no significant interaction was observed, when an egocentric strategy was present, adding an exocentric component increased likability. This pattern implies that likability is shaped not simply by the absolute level of altruism but also by the context of interactions, including participants’ expectations. An agent that largely maintains its own policy while occasionally

accommodating the participant's interests may be interpreted as having a balanced or negotiable stance, which can enhance positive evaluation.

Finally, the perceived intelligence was lower for agents employing egocentric strategies. Intuitively, one may expect agents who consistently follow their own policies to evoke perceptions of greater strength of will, shrewdness, and other intellectual abilities, whereas agents who merely react to user feedback do not. However, the present results show the opposite pattern, thus highlighting the importance of cognition based on social context, similar to the results of the likability.

In social context, intelligence is often evaluated based on an individual's ability to act wisely in human relations (Thorndike, 1920). In negotiation settings like the ultimatum game, rejection functions as a critical signal communicating that a proposal is perceived as unjust, and participants inherently expect proposers to recognize and respond to this feedback. Therefore, the egocentric agents' failure to adapt was likely perceived as a fundamental lack of situational awareness—an inability to understand the social context of the interaction.

For an agent's persistence to be evaluated as “determined” rather than merely “stubborn,” the observer may need to be able to layer a goal-oriented motive onto the observed rigidity. In the present study, however, the egocentric agents followed a probabilistic policy that lacked a clear objective. This lack of socially interpretable cues likely prevented participants from upgrading their interpretation from “stubbornness” to “determination.” This is also consistent with previous data showing that an “egocentric agent”—bearing the same label but, unlike the present egocentric agent, consistently attempting to behave in ways advantageous to itself—was evaluated as more intelligent than the other conditions (Oasa et al., 2026). Thus, perceived intelligence may reflect not the internal complexity of an agent's policy but the degree to which its behavior aligns with socially expected patterns and supports an interpretable intention.

The dissociation in subjective ratings—where exocentric strategies increased perceived safety while egocentric strategies decreased perceived intelligence—indicates that participants' representations of others are multidimensional rather than reducible to a single evaluative dimension such as the amount of distribution; if participants had reacted solely to the reward, their evaluations would likely have shifted uniformly across all traits. From a computational perspective, this suggests that participants infer multiple latent properties of the agent's policy, such as responsiveness to feedback, rigidity of internal constraints, and overall cooperativeness.

Taken together, these results support the view that human other-models incorporate representations of action-generation principles, rather than relying solely on outcome-based evaluations. Future work can test this mechanism more directly by modeling accept–reject decisions at the trial level using learning models that update beliefs about an agent's responsiveness and rigidity based on observed behavioral transitions.

Summary

In this study, we investigated the emergence of reciprocal altruism using an ultimatum game between human participants and artificial agents with different behavioral strategies.

The results align with reciprocal altruism theories, which posit that individuals cooperate with partners from whom they expect reciprocity and sanction those who do not. Beyond outcome-level effects, behavioral and subjective measures suggest that participants formed an other-model incorporating process cues such as responsiveness and policy rigidity.

Overall, this study shows that humans are sensitive not only to how much others give but also to how actions are generated in response to social feedback, with implications for fairness-related decision-making and impression formation.

However, several limitations remain. First, the extent to which the observed interaction patterns with artificial agents generalize to interactions between two human players is unclear; future work should compare agent-based and fully human–human interactions under matched conditions. An important next step is to develop a decision-making model based on the insights from the present study and quantitatively compare its predictions with human performance.

Second, the agent algorithms can be further refined. In particular, the parameters used in the Q-learning model for the exocentric/egocentric condition were fixed a priori and thus should be calibrated based on empirical human data.

Third, as the points earned in the game did not correspond directly to actual monetary incentives, the participants' economic decisions may not fully reflect real-world stakes. However, rejection of unfair offers in the ultimatum game has been shown to emerge robustly even when the absolute amount of money at stake varies (Cameron, 1999), and meta-analytic evidence suggests that monetary incentives do not fundamentally alter the qualitative pattern of behavior in interactive decision-making tasks, including the ultimatum game (Camerer & Hogarth, 1999). Moreover, because the primary aim of the present study was to examine how people infer the latent policies of agents rather than to estimate precise monetary valuations, this limitation is unlikely to substantially undermine the validity of the main conclusions.

Finally, the present design does not fully disentangle whether participants reacted to the strategic responsiveness per se or to the resulting offer distribution. Future studies should include explicit manipulation checks (e.g., evaluate perceived responsiveness/fairness using open-ended questionnaire) and control offer distributions to isolate process-level inferences. One promising approach would be to introduce a yoked-control agent condition, in which the round-by-round offers generated in an actual interaction with an agent are recorded and then replayed in the same order to a different participant.

Addressing these issues through refined experimental control will further deepen our understanding of the dynamic mechanisms underlying social interaction.

References

- Bartneck, C. (2023). Godspeed questionnaire series: Translations and usage. In *International handbook of behavioral health assessment* (pp. 1–35). Springer. https://doi.org/10.1007/978-3-030-89738-3_24-1
- Camerer, C. F., & Hogarth, R. M. (1999). The effects of financial incentives in experiments: A review and capital-labor-production framework. *Journal of Risk and Uncertainty*, 19(1/3), 7–42. <https://doi.org/10.1023/A:1007850605129>
- Cameron, L. A. (1999). Raising the stakes in the ultimatum game: Experimental evidence from indonesia. *Economic Inquiry*, 37(1), 47–59.
- Falk, A., & Fischbacher, U. (2006). A theory of reciprocity. *Games and Economic Behavior*, 54(2), 293–315. <https://doi.org/https://doi.org/10.1016/j.geb.2005.03.001>
- Fehr, E., & Fischbacher, U. (2003). The nature of human altruism. *Nature*, 425, 785–91. <https://doi.org/10.1038/nature02043>
- Fehr, E., & Gächter, S. (2002). Altruistic punishment in humans. *Nature*, 415, 137–40. <https://doi.org/10.1038/415137a>
- Fehr, E., & Schmidt, K. M. (1999). A theory of fairness, competition, and cooperation. *The Quarterly Journal of Economics*, 114(3), 817–868. <https://doi.org/10.2139/ssrn.106228>
- Güth, W., Schmittberger, R., & Schwarze, B. (1982). An experimental analysis of ultimatum bargaining. *Journal of Economic Behavior & Organization*, 3(4), 367–388. [https://doi.org/10.1016/0167-2681\(82\)90011-7](https://doi.org/10.1016/0167-2681(82)90011-7)
- György, G., & Watson, J. (1999). Early socio-emotional development: Contingency perception and the social-biofeedback model. *Early Social Cognition*, 101–136.
- Hamada, D., Otsu, K., & Hayashi, Y. (2023). Behavioral characteristics in general trust: An exploratory laboratory-based analysis using the ultimatum game. *Proceedings of the 45th Annual Conference of the Cognitive Science Society*, 1815–1822.
- Hamilton, W. (1964). The genetical evolution of social behaviour. ii. *Journal of Theoretical Biology*, 7(1), 17–52. [https://doi.org/10.1016/0022-5193\(64\)90039-6](https://doi.org/10.1016/0022-5193(64)90039-6)
- Hayashi, Y., & Okada, R. (2017). Compound effects of expectations and actual behaviors in human-agent interaction: Experimental investigation using the ultimatum game. *Proceedings of the 39th Annual Meeting of the Cognitive Science Society*, 2168–2173.
- Hoffman, G., & Breazeal, C. (2007). Cost-based anticipatory action selection for human–robot fluency. *IEEE Transactions on Robotics*, 23(5), 952–961. <https://doi.org/10.1109/TRO.2007.907483>
- Johnson, S. C. (2000). The recognition of mentalistic agents in infancy. *Trends in Cognitive Sciences*, 4(1), 22–28. [https://doi.org/https://doi.org/10.1016/S1364-6613\(99\)01414-X](https://doi.org/https://doi.org/10.1016/S1364-6613(99)01414-X)
- Kasaneko. (2019, November). Gabriel (in japanese). <https://hub.vroid.com/characters/7011619231874599278/models/7613529040682918185>
- Kusatake, D. (2020). *Experimental investigation of the ultimatum game in vr environment (in japanese)* [Informally published undergraduate thesis, College of Comprehensive Psychology, Ritsumeikan University, Osaka, Japan].
- Mochizuki, F. (2014). On the response to "altruistic" offers in ultimatum games (in japanese). *Studies in moralogy : the forum for ethical and moral studies*, (73), 25–39.
- Nass, C., Fogg, B. J., & Moon, Y. (1996). Can computers be teammates? *International Journal of Human-Computer Studies*, 45(6), 669–678. <https://doi.org/10.1006/ijhc.1996.0073>
- Nowak, M. A., & Sigmund, K. (1992). Tit for tat in heterogeneous populations. *Nature*, 355(6357), 250–253. <https://doi.org/10.1038/355250a0>
- Oasa, K., Shimojo, S., & Hayashi, Y. (2026). Influence of agent's strategy on individual's cognition and decision making: Experimental investigation using ultimatum game. *Proceedings of the 13th International Conference on Human-Agent Interaction*, 222–228. <https://doi.org/10.1145/3765766.3765794>
- Premack, D., & Woodruff, G. (1978). Does a chimpanzee have a theory of mind. *Behavioral and Brain Sciences*, 1, 515–526. <https://doi.org/10.1017/S0140525X00076512>
- Rabin, M. (1993). Incorporating fairness into game theory and economics. *The American Economic Review*, 83(5), 1281–1302.
- Stephens, C. (1996). Modelling reciprocal altruism. *The British Journal for the Philosophy of Science*, 47(4), 533–551. <https://doi.org/10.1093%2Fbjps%2F47.4.533>
- Sutter, M. (2007). Outcomes versus intentions: On the nature of fair behavior and its development with age. *Journal of Economic Psychology*, 28(1), 69–78. <https://doi.org/10.1016/j.joep.2006.09.001>
- Thorndike, E. L. (1920). Intelligence and its uses. *Harper's Magazine*, 140, 227–235.
- Trivers, R. (1971). The evolution of reciprocal altruism. *Quarterly Review of Biology*, 46, 35–57. <https://doi.org/10.1086/406755>
- Watson, J. S. (1972). Smiling, cooing, and "the game." *Merrill-Palmer Quarterly of Behavior and Development*, 18(4), 323–339.
- Yoshida, W., Dolan, R., & Friston, K. (2009). Game theory of mind. *PLoS computational biology*, 4, e1000254. <https://doi.org/10.1371/journal.pcbi.1000254>