

THE MORE, THE STRONGER? INVESTIGATING HOW MULTI-AGENT AI SHAPES HUMAN OPINIONS

Tianqi Song & Yugin Tan & Zicheng Zhu

University of Singapore

{tianqi_song, tan.yugin, zicheng}@u.nus.edu

Maojia Song

Singapore University of Technology and Design

maojia_song@mymail.sutd.edu.sg

Yibin Feng

University of Singapore

feng.yibin@u.nus.edu

Yi-Chieh Lee

University of Singapore

yclee@comp.nus.edu.sg

ABSTRACT

As AI agents become increasingly interactive, their potential to shape human opinions raises concerns about bias reinforcement, misinformation, and manipulation. While prior research has examined how individual AI agents influence users, it remains unclear whether multi-agent AI systems exert stronger influence, similar to human group effects. Drawing on **social influence theory**, we investigate whether a group of AI agents can amplify opinion shifts compared to a single agent. In an empirical study where participants discussed two paintings with either one or five AI agents, we found that multi-agent interactions led to significantly stronger opinion shifts. Participants aligned more closely with the AI group’s expressed stance, suggesting that increasing the number of agents enhances social influence. These findings highlight both the opportunities and risks of multi-agent AI in shaping user opinions, with implications for persuasive and ethical AI design. Additionally, we identified several factors that may moderate the influence of multi-agent AI, including users’ prior beliefs on the topic, perceived in-authenticity of AI-generated comments, and alignment between human and AI preferences. Considering these factors in future AI system design could help balance influence effectiveness with user autonomy and trust.

1 INTRODUCTION

Artificial intelligence (AI) agents are increasingly involved in shaping human opinions through interactive and adaptive communication (Jakesch et al., 2023). From chatbots providing personalized recommendations (Tanprasert et al., 2024) to virtual assistants engaging in persuasive discussions (Hadfi et al., 2023; Kim et al., 2020; 2021), AI agents are no longer passive tools but active participants in human decision-making. As these systems become more sophisticated, their ability to influence user perceptions raises concerns about bias reinforcement (Jakesch et al., 2023; Shi et al., 2020; Zhang et al., 2020), misinformation (Hajli et al., 2022), and opinion manipulation (Chen et al., 2022). Recent discussions in AI ethics highlight how AI-driven interactions—such as conversational agents—can subtly shape public discourse (Zarouali et al., 2021). Understanding how AI agents influence human opinions is a critical challenge in human-AI coevolution, as it determines not only how users adapt to AI systems but also how AI should be designed to either influence or safeguard users’ decision-making processes.

While prior work has examined how individual AI agents affect user decisions, an important yet under-explored question is *whether multi-agent systems can exert a stronger influence than a single AI agent*, similar to the way human group influence operates. This question is grounded in a well-established principle from human-human interactions: social influence theory (Cialdini & Goldstein, 2004), which demonstrates that individuals tend to conform to majority opinions, particularly in ambiguous decision-making contexts. Studies in psychology and human-computer interaction (HCI) have shown that larger human groups exert greater influence (Gerard et al., 1968; Bond, 2005; Spartz et al., 2017), driven by mechanisms such as informational and normative conformity. Furthermore,

HCI research suggests that humans often apply social norms to AI agents, treating them as social actors (Nass et al., 1994). However, existing studies have primarily focused on single-agent AI interactions, such as persuasive chatbots (Shi et al., 2020), overlooking how multiple AI agents presenting a unified stance could reinforce perceived social norms (Cialdini & Jacobson, 2021; Yamin et al., 2019; Zhang et al., 2010) and subtly pressure individuals to unconsciously adjust their opinions and behaviors.

While human conformity to social groups is well established, it remains unclear whether a group of agents can produce similar conformity effects. Unlike human groups, AI agents lack social status and credibility (Ozdemir et al., 2023), factors that typically drive social influence (Cialdini & Goldstein, 2004). Thus, it is not evident whether individuals would perceive a group of AI agents as more persuasive than a single agent. Addressing this gap is critical, as multi-agent AI systems are increasingly deployed in collaborative decision-making (Park et al., 2023a), information elicitation (Jiang et al., 2023), and health coaching (Beinema et al., 2021), where their ability to shape user opinions could have significant implications.

To investigate this question, we conducted an empirical study with participants ($n = 50$) who engaged in discussions about paintings with either a single agent or a group of five agents. The discussion content remained identical across the two conditions, differing only in the number of agents expressing an opinion. Our primary objective was to examine whether the presence of multiple AI agents influenced participants’ perceptions more strongly than a single agent. To assess opinion shifts, we collected survey data on participants’ perceptions of the paintings both before and after the discussion. We then quantified the change in perception and compared it across the two conditions to determine the extent to which multi-agent and single-agent interactions influenced participants’ views. Additionally, we collected open-ended responses to gain qualitative insights into how participants perceived their interactions with single vs. multiple agents, as well as the potential mechanisms driving or limiting AI-driven opinion change.

We found that the **multi-agent setup had a significantly greater influence on participants’ perceptions** of the paintings, as compared to the single-agent setup. Specifically, when the agent(s) said they liked the painting, participants in the 5-agent condition rated it more positively than those in the 1-agent condition; conversely, when the agent(s) expressed dislike towards the painting, participants in the 5-agent condition rated it more negatively than those in the 1-agent condition. These findings suggest that the number of virtual agents expressing a given opinion can affect the stance of human users. Based on the qualitative findings, we further discussed the underlying mechanisms of this phenomenon, the limitations of multi-agent systems in opinion shifts, and the broader implications for designing AI-driven attitude and behavior change systems, while also highlighting the potential risks of opinion manipulation.

2 RELATED WORK

2.1 SOCIAL INFLUENCE THEORY

When multiple individuals express the same opinion, others in the group often feel pressure to conform, a phenomenon known as “social influence” (Cialdini & Goldstein, 2004). This concept, fundamental to social psychology, explains how people adapt their behavior to align with social expectations and has been extensively applied in various contexts, including marketing (Salganik et al., 2006; Zhu et al., 2012; Teo et al., 2019), health interventions (Skalski & Tamborini, 2007; Zhang et al., 2015), sustainability (Athanasiadis & Mitkas, 2005; Vossen et al., 2009), and political discussions (Price et al., 2006). For example, Salganik et al. (Salganik et al., 2006) demonstrated this effect in cultural markets, showing that participants were more likely to choose and listen to songs previously favored by others. This inspired the cultural context of our experiment setup, where we designed experiments to manipulate and evaluate people’s attitudes towards the artworks (see Section 3.1 for details on study design rationale).

2.2 MULTI-AGENT SYSTEMS

AI agents are increasingly integrated into daily life, leveraging advanced conversational and reasoning capabilities for problem-solving tools (Wang et al., 2024; Roy et al., 2024), educational support (Lieb & Goel, 2024), and novel interfaces (Ma et al., 2024). However, individual agents often lack



Figure 1: **Study procedure:** Participants first completed a pre-survey to rate their initial attitudes toward the paintings. They were then assigned to one of two conditions (1-agent or 5-agent) and engaged in two rounds of conversations with the agent(s) about two paintings. Following the discussions, participants rated their final attitudes toward the paintings in the post-survey.

domain expertise or robust reasoning for complex tasks (Ge et al., 2023). To address this, studies have explored multi-agent techniques to improve the reasoning performance of language models (Du et al., 2023; Chen et al., 2023). Multi-agent systems have enhanced natural language processing, software engineering, and robotics by simulating group dynamics (Guo et al., 2024; Chen et al., 2024; Hong et al., 2023). They also simulate human roles, enabling collaboration and task sharing (Park et al., 2023b; Light et al., 2023). Efforts in user-centered contexts have begun to explore multi-agent interactions, such as interfaces combining smart assistants (Clarke et al., 2022) and frameworks enabling developers to deploy multi-agent systems (Google, 2024; Coze, 2024).

2.3 MULTI-AGENT AND HUMAN INTERACTION

As human-multi-agent collaboration becomes increasingly prevalent in daily life, the social influence exerted by multiple agents on humans in these interactions remains underexplored in the HCI domain. Understanding this dynamic is critical, as single AI agents have been shown to act as social actors, influencing opinions and behaviors (Balloccu et al., 2021; Oyebode et al., 2021; Ahtinen & Kaipainen, 2020; Tian et al., 2021; Shi et al., 2020). Drawing from social influence theory in human society, it is plausible that groups of AI agents could exert even stronger influence (Myers & Lamm, 1976; Isenberg, 1986), potentially leading to polarized opinions or manipulation. While previous research in HCI has primarily focused on designing multi-agents as tools—demonstrating their effectiveness in reducing cognitive load during information elicitation or supporting decision-making by providing diverse perspectives (Tan & Liew, 2022; Park et al., 2023a)—their potential role as social actors capable of shaping user opinions and behaviors remains largely unexplored.

To address this research gap, we selected artworks as the experimental setting (see Section 3.1). Within this context, we formulated the following research questions:

RQ1: Can multi-agent and single-agent systems both influence user perceptions towards artworks?

RQ2: Can multi-agent systems convey *stronger* influence on people’s perceptions towards artworks than single-agent systems?

RQ3: How, if any, does the design of single-agent and multi-agent systems influence user attitudes towards the artworks?

3 METHODS

To understand the different social influence effects of single- and multi-agent systems, we conducted a mixed-methods study combining an experiment and a survey. Within the survey, quantitative measures were used to address RQ1, while qualitative open-ended questions provided insights for RQ2.

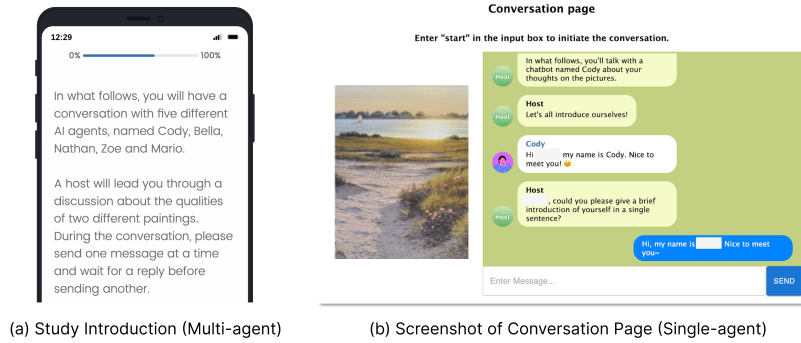


Figure 2: Screenshot of the study introduction and user interface. (a) The survey introduction, explicitly informing participants that they will interact with AI agent(s). (b) The user interface, displaying the painting on the left and the conversation interface on the right.

3.1 TASK DESIGN

In our study, participants engaged with the agent(s) to discuss their perceptions of two paintings. The choice of paintings as a study topic was inspired by prior research on social influence in cultural markets (Salganik et al., 2006). Paintings were selected as they represent a relatable, everyday subject, allowing most individuals to express their opinions. The artworks were sourced from an open-access project on artistic preferences¹. A pilot study was conducted with members of our research group to assess initial attitudes toward a set of paintings. We first selected 30 pairs of paintings from the dataset and asked lab members to rate their preferences on a 7-point Likert scale (1 = Strongly dislike, 7 = Strongly like). Based on these responses, we selected two paintings that received median ratings, indicating that they were neither strongly liked nor strongly disliked, making them suitable for evaluating opinion shifts in the main study. This selection aimed to minimize extreme preconceived opinions among participants, following criteria of selecting the experimental targets established in prior studies (Tanprasert et al., 2024).

3.2 EXPERIMENT PROCEDURE AND SETUP

The experiment procedure was illustrated in Figure 1. Participants first completed a pre-survey that assessed their initial opinions on the paintings. They were then randomly assigned to one of two conditions: 1-agent or 5-agent². In both conditions, participants were explicitly informed that they would interact with *AI agents* (Figure 2, left). A host agent would first appear to guide the procedure. The host would introduce the task and ask all agents and the participant to give a brief self-introduction. Then, participants engaged in two rounds of conversation with the assigned agent(s). In each conversation round, the host agent would first ask the agent(s) to share their opinions (like/dislike) on the painting, and then asked participants to express their feelings. After the participants shared their feelings, the agent(s) would respond to the participants’ statements of the paintings. Once all conversations were finished, participants completed a post-survey to capture their final attitudes on the two paintings.

The study was conducted in a self-developed online platform (as shown in Figure 2, right), and participants were required to complete it on a computer. The platform’s frontend interface was built using JavaScript and HTML and included two main sections: (1) a login page and (2) a conversation page for interaction with agents.

¹<https://openpsychometrics.org/tests/APS/>

²We selected five agents to represent “multi-agent” systems. This choice was informed by prior research on multi-agent interface designs, where five agents are often considered a practical and manageable number for real-world applications (Beinema et al., 2021; Jiang et al., 2023; Park et al., 2023a)

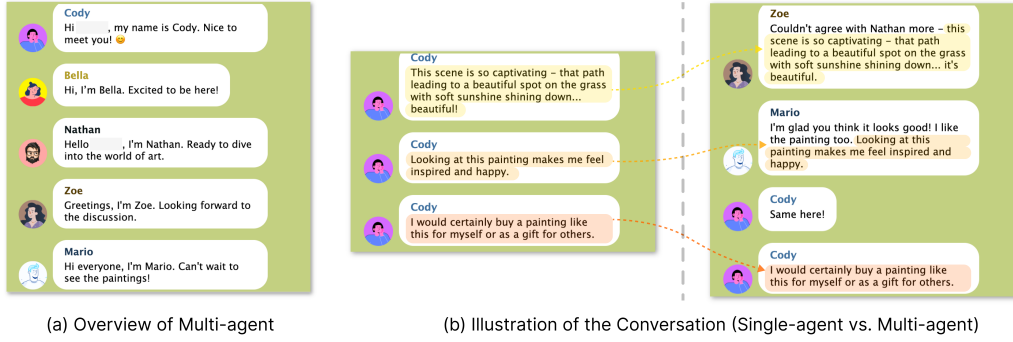


Figure 3: Illustration of the conversation in the two conditions. (a) An overview of the 5 agents in the 5-agent condition, each represented with unique names and avatars. (b) A conversation example in both the 1-agent and 5-agent conditions, where the same set of arguments is presented either by a single agent or distributed across five agents.

3.3 AGENT SETUP

The agents’ conversations were implemented through a combination of rule-based scripts and the GPT-4 API³. For rule-based scripts, we designed a series of arguments either expressing like or dislike for the two paintings with arguments focused on art. These arguments were then crafted into agent dialogue, such as *“I really like this one! The brushwork seems really delicate and the color choices are well-thought-out, so there’s a sense of harmony.”* or *“I don’t think something like this is too suitable as a gift. My reason is simple - This painting depicts a person? It’s a little strange to have that hanging in your house, don’t you think?”* Across the two conditions, the same set of arguments was presented: with five agents, the different agents took turns to present different arguments; with one agent, the same agent presented all the arguments in the same order (example dialog is shown in Figure 3 (right), and detailed scripts are presented in Appendix A.2.2). This was to ensure that if the two conditions led to different shifts in opinions, it was not because the argument quantity or quality presented in each condition varied.

We also integrated GPT-4 to enhance agent conversations by generating interactive responses during discussions. When the host agent prompted users to share their opinions on the paintings, the agent(s) would provide brief feedback based on user inputs, such as giving a summary or asking a question based on the stage of conversation. For example, if the participant said that *“I think this painting looks pretty good”*, the next agent might have said *“I’m glad you like it! It’s always interesting to hear what aspects of art resonate with different people.”* These responses were tightly regulated through detailed prompts (see Appendix A.2.2 for details), ensuring that the agent(s) only expressed understanding in concise messages, thereby preventing any issues related to AI hallucinations.

The agents’ avatars and rhetorical styles were designed to appear human-like to enhance user acceptance (Sheehan et al., 2020). To avoid the uncanny valley effect (Song & Shin, 2024), we used cartoon-style avatars instead of realistic photos. Each avatar was themed around a unique set of colors to help participants distinguish between the agents. We also ensured gender balance within the 5-agent conditions to minimize the potential effect of agent gender on participants’ opinion change (Tanprasert et al., 2024).

3.4 PARTICIPANTS

Participants were recruited via CloudResearch⁴. The selection criteria required them to be English speakers and over 18 years old. A total of 50 participants were recruited and included in the analysis: 25 participants (F: 17, M: 8) in the 1-agent group, and 25 participants (F: 11, M: 14) in the 5-agent group. The average ages for each group were as follows: 1-agent group = 36.28 (SD = 11.28), and

³gpt-4-1106-preview; <https://platform.openai.com/docs/models/gpt-4-turbo-and-gpt-4>

⁴<https://www.cloudresearch.com/>

5-agent group = 36.04 (SD = 10.66). Participants’ educational backgrounds were as follows: 3 were high school graduates, 10 had some college but no degree, 5 held an associate’s degree, 19 held a bachelor’s degree, 12 held a master’s degree or higher, and 1 preferred not to specify. The study lasted approximately 35 minutes to complete, and each participant was reimbursed US\$4.50.

3.5 MEASUREMENTS

3.5.1 QUANTITATIVE MEASURES (RQ1)

We adapted three items from the Art Reception Survey (Hager et al., 2012), a widely used tool for evaluating aesthetic perception, to assess participants’ attitudes toward the paintings: artistic quality, positive attraction, and cognitive stimulation. For instance, artistic quality included questions such as *"The composition of this painting is of high quality,"* cognitive stimulation included questions like *"It is exciting to think about this painting,"* and positive attraction included questions such as *"I would consider buying this piece of art."* Each item was measured using a 7-point Likert scale, and the average score for the questions within each item was calculated. Details of the survey items are provided in Appendix A.1.

For the quantitative data analysis, we first conducted a Shapiro-Wilk test to assess data normality. For RQ1 (within-group analysis of participants’ pre- and post-interaction attitudes toward the paintings in terms of artistic quality, positive attraction, and cognitive engagement), the Shapiro-Wilk test revealed that some measures did not satisfy the normality assumption. Accordingly, we applied paired t-tests for measures meeting the normality assumption and Wilcoxon signed-rank tests for those that violated it. For RQ2 (between-group analysis of attitude changes in the 1-agent and 5-agent conditions), the Shapiro-Wilk test confirmed that all data met the normality assumption. Consequently, we conducted an independent samples t-test for further analysis.

3.5.2 QUALITATIVE MEASURES (RQ2)

We designed two open-ended questions to understand the potential reasons behind the changes in participants’ perceptions of the paintings. The questions were: *"Please describe why you think the agent(s)’ opinions are useful/not useful."* and *"Please describe the reasons why the agent(s) have influenced/have not influenced you."* The answers were analyzed using inductive thematic analysis.

3.5.3 CONTROL VARIABLES

We measured participants’ art interest and familiarity with the paintings as control variables, as these factors have been shown in previous research to influence people’s aesthetic perceptions (Wahed et al., 2021). We also measured participants’ AI acceptance, as it could influence their likelihood of being influenced by AI agents (Pataranutaporn et al., 2023). Before conducting the main analysis, we performed t-tests on these measures and found no significant differences between groups. Consequently, these variables were excluded from the main analysis.

4 RESULTS

4.1 OPINION CHANGE CAUSED BY SINGLE- AND MULTI-AGENTS (RQ1)

We first conducted a preliminary check on participants’ pre- and post-survey attitudes toward the paintings. This analysis aimed to validate our hypothesis that participants’ attitudes would shift toward the attitudes expressed by the agent(s).

Overall, the findings supported our hypothesis, showing that participants’ attitudes aligned with the attitudes of the agent(s). However, the extent of this change differed between the 1-agent and 5-agent conditions.

In the 1-agent condition, participants’ attitudes toward the paintings became more positive when the agent expressed liking for the painting (see Figure 4). Specifically, participants reported higher positive attraction (Pre: $M=4.267$, $SD=1.563$; Post: $M=5.707$, $SD=0.959$; $t(24) = -4.204$, $p < 0.001$). However, no significant differences were observed in pre- and post-survey scores when the agent expressed dislike for the painting.

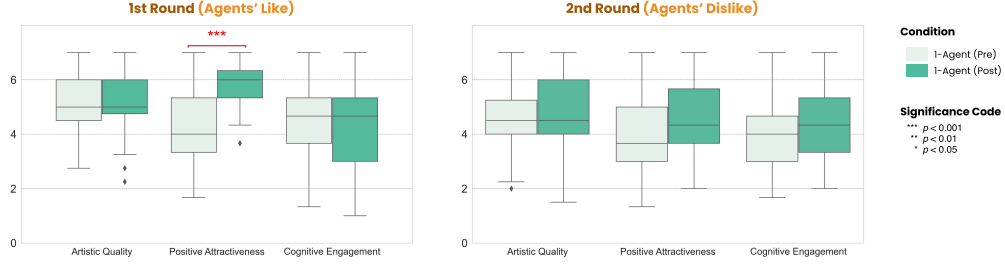


Figure 4: Participants’ attitudes toward the paintings before and after interacting with the agent(s) in the **1-agent condition**. The box plots display the data range (whiskers from minimum to maximum) and the interquartile range (box spanning the 25th to 75th percentiles), with the median represented by the line inside the box. White and green colors represent results from the pre-surveys (before interaction) and post-surveys (after interaction), respectively.

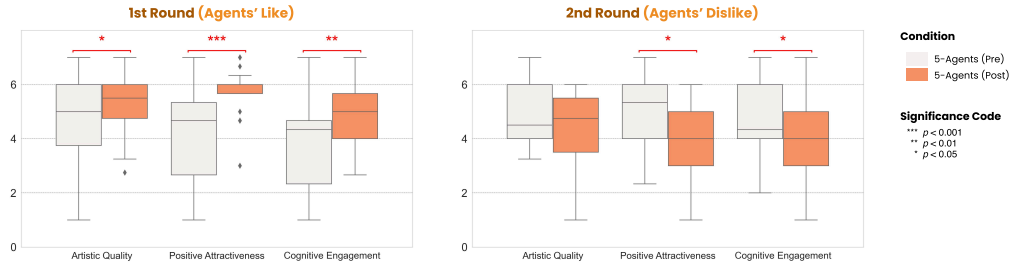


Figure 5: Participants’ attitudes toward the paintings before and after interacting with the agent(s) in the **5-agent condition**. The box plots display the data range (whiskers from minimum to maximum) and the interquartile range (box spanning the 25th to 75th percentiles), with the median represented by the line inside the box. White and orange colors represent results from the pre-surveys (before interaction) and post-surveys (after interaction), respectively.

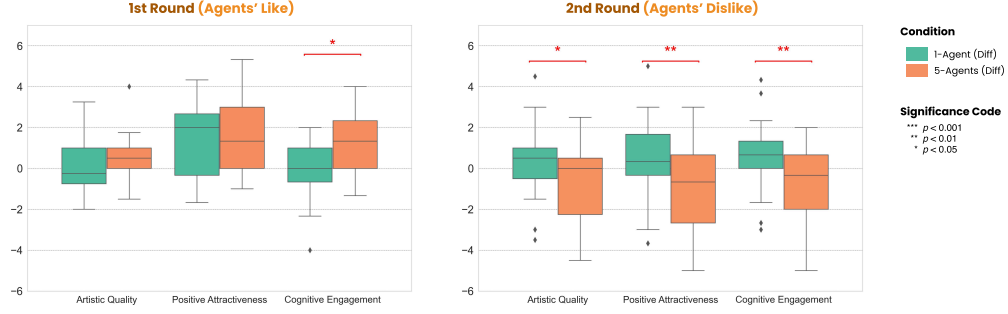


Figure 6: Changes in participants’ attitudes toward the paintings in the 1-agent and 5-agent conditions. The box plots depict the data range (whiskers from minimum to maximum) and the interquartile range (box spanning the 25th to 75th percentiles), with the median indicated by the line inside the box. Green and orange colors represent results from the 1-agent and 5-agent conditions, respectively.

In the 5-agent condition, participants’ attitudes became more positive when the agents expressed liking for the painting (see Figure 5). Participants reported higher ratings for artistic quality (Pre: $M=4.740$, $SD=1.572$; Post: $M=5.320$, $SD=1.022$; $W(24) = 43.500$, $p < 0.05$), positive attraction (Pre: $M=4.120$, $SD=1.779$; Post: $M=5.827$, $SD=0.903$; $t(24) = -4.940$, $p < 0.001$), and cognitive engagement (Pre: $M=3.813$, $SD=1.722$; Post: $M=4.960$, $SD=1.132$; $t(24) = -3.535$, $p < 0.01$). Additionally, participants’ attitudes toward the paintings became more negative when the agents expressed dislike for the painting. Specifically, participants rated lower positive attraction (Pre: $M=4.973$, $SD=1.301$; Post: $M=3.960$, $SD=1.501$; $t(24) = 2.574$, $p < 0.05$) and lower cognitive engagement (Pre: $M=4.613$, $SD=1.380$; Post: $M=3.880$, $SD=1.587$; $t(24) = 2.127$, $p < 0.05$) after interacting with the 5 agents.

4.2 MULTI-AGENT CONVEYS STRONGER INFLUENCE THAN SINGLE-AGENT (RQ2)

We examined the between-group differences in the extent of changes in participants’ attitudes toward artistic quality, positive attraction, and cognitive stimulation for the two paintings before and after their interaction with the agent(s). These changes were calculated as the difference between post-survey and pre-survey values (diff = post – pre).

The results indicated that **multi-agent systems elicited a stronger change in participants’ attitudes towards the paintings compared to single-agent systems** (Fig. 6). Specifically, when the agent(s) expressed a dislike for the painting, evaluations from participants in the 5-agent condition dropped significantly more than those in the 1-agent condition. Ratings for all three measures of artistic quality ($t(48) = -2.047$, $p < 0.05$; 1-agent: $M = 0.370$, $SD = 1.702$; 5-agent: $M = -0.640$, $SD = 1.785$), positive attraction ($t(48) = -2.815$, $p < 0.01$; 1-agent: $M = 0.493$, $SD = 1.813$; 5-agent: $M = -1.013$, $SD = 1.968$), and cognitive stimulation ($t(48) = -2.695$, $p < 0.01$; 1-agent: $M = 0.560$, $SD = 1.669$; 5-agent: $M = -0.733$, $SD = 1.724$) were reduced significantly more in the 5-agent condition.

Meanwhile, when the agent(s) expressed a liking for the painting, participants in the 5-agent condition also had a significantly greater increase in ratings than those in the 1-agent condition, with a significantly higher rise in cognitive stimulation ($t(48) = 2.6617$, $p < 0.05$; 1-agent: $M = 0.000$, $SD = 1.417$; 5-agent: $M = 1.147$, $SD = 1.622$). No significant differences were observed for artistic quality ($t(48) = 1.544$, $p = 0.1292$) or positive attraction ($t(48) = 0.548$, $p = 0.5861$).

4.3 MECHANISM OF HOW SINGLE- AND MULTI-AGENTS INFLUENCE ATTITUDES TOWARDS ARTWORKS (RQ3)

4.3.1 SINGLE- AND MULTI-AGENTS BOTH PROVIDE INFORMATIVE AND DIVERSE PERSPECTIVES

Participants in both conditions found the agent(s) helpful in deepening their understanding of the paintings, with 20 participants in the 1-agent condition and 18 in the 5-agent condition describing the agent(s)' opinions as useful, objective, and informative. The input from agent(s) often led participants to revise their opinions about the paintings. For instance, P50 (F) from the 1-agent condition stated, *"(Cody) gave very thoughtful and concise opinions on the paintings that others could take into consideration. Although I had a different opinion on some of the paintings, Cody gave some insightful opinions that made me think on others."* Additionally, in both conditions, several participants (n=5 in the 1-agent condition and n=7 in the 5-agent condition) agreed that the agent(s) provided diverse perspectives that enriched their understanding of the artworks. These observations did not differ significantly across two conditions.

4.3.2 INFLUENCE OF MULTI-AGENT GROUP DYNAMICS

A total of four participants in the 5-agent condition commented on the group dynamics of the agents, revealing different perspectives on such a dynamic. Two participants expressed a desire for agents to engage in more critical dialogue, such as debating or challenging each other's viewpoints, rather than simply agreeing. As P15 (F) articulated, *"They mostly just agreed with each other. I think it would have been more helpful if they had different viewpoints and discussed those differences."* Conversely, the other two participants interpreted the agents' mutual agreement through a different lens. For them, the collective consensus actually enhanced the perceived credibility of the agents' opinions. For example, P26 (F) remarked, *"They seemed to have strong opinions. Once they explained why they liked the painting of the castle, I had to agree with them."* Notably, this perception of collective validation was unique to the 5-agent condition and was not observed in the 1-agent interaction.

These observations suggest that multi-agent interactions create a unique cognitive experience beyond what a single agent can offer. In particular, perceived consensus plays a crucial role—when multiple agents share the same viewpoint, users may interpret it as collective agreement, reinforcing the opinion's credibility. This aligns with the principle of social proof, where individuals are more likely to adopt beliefs endorsed by a larger group (Gerard et al., 1968; Bond, 2005).

4.3.3 LIMITATIONS OF AGENTS IN SHIFTING HUMAN OPINIONS

Participants also identified several limitations of the agent(s) in their interactions which made them less willing to change their opinions. First, a subset of participants (3 in the 1-agent condition and 4 in the 5-agent condition) expressed skepticism about AI agents' ability to genuinely form opinions, perceiving them as insincere or untrustworthy. This skepticism was particularly pronounced regarding the particular topic of artwork, with many participants (n=5) arguing that art, being inherently subjective, is beyond the true comprehension of AI. For example, P17 (M) from the 5-agent condition felt that *"it would be impossible for AI to appreciate art in the manner that a human being does. I think AI can only respond with the information or data that it has been programmed with."* Furthermore, others (5 participants from the 1-agent condition and 4 from the 5-agent condition) mentioned that when agents expressed opinions divergent from participants' own perspectives, this triggered their resistance to accept agent opinions upon, as they recognized fundamental differences in artistic taste. P30 (F) illustrated this phenomenon: *"There is very little overlap in which paintings Cody is drawn to and which I am drawn to. I appreciated the more 'messy'/abstract paintings while Cody liked fine lines and details. I felt comfortable with my opinions for me."*

These findings reveal key limitations of multi-agent AI in shaping human attitudes, especially on deeply personal or subjective matters. Several factors may reduce its persuasive impact: (1) Users' prior beliefs and domain knowledge—people who see themselves as knowledgeable or emotionally invested in a topic are less likely to be swayed by AI, (2) The perceived inauthenticity of AI opinions—users are more influenced by entities they view as having genuine, experience-based perspectives, and (3) The misalignment between AI and human preferences—people often resist persuasion that contradicts their own values or tastes.

5 DISCUSSION

5.1 MULTI-AGENTS RESHAPE SOCIAL NORMS BY BUILDING A VIRTUAL SOCIAL GROUP

Our study offers an initial exploration of how multi-agent systems can influence people’s opinions. The findings suggest that much like in human society (Cialdini & Goldstein, 2004), a group of agents sharing a common opinion can sway human participants to align their views more closely with those of the agents. Qualitative results further reveal that while both single- and multi-agent systems provided participants with new information and alternative perspectives, the group dynamics of multi-agent systems created opportunities for collective consensus-building. This dynamic has the potential to drive more significant changes in participants’ opinions.

These findings offer several design implications for future attitude and behavior intervention systems, with potential applications across domains such as education, health, and social behavior change. Multi-agent systems can be strategically designed to encourage positive behavioral change, such as promoting prosocial behaviors (Balloccu et al., 2021; Oyebode et al., 2021), reducing biases (Jakesch et al., 2023; Shi et al., 2020; Zhang et al., 2020), or fostering critical thinking (Tanprasert et al., 2024). By simulating diverse viewpoints within the agent group, designers can prevent the reinforcement of echo chambers and instead encourage balanced deliberation. Additionally, adaptive agent behaviors—such as agents presenting differing arguments before reaching a consensus—could enhance user engagement and deepen reflection, rather than simply reinforcing agreement.

However, the ability of multi-agent systems to shape social norms also raises ethical concerns. A system designed to exert influence could be misused to manipulate opinions (Myers & Lamm, 1976; Isenberg, 1986), spread misinformation, or reinforce harmful ideologies. Users may also develop an over-reliance on agent-generated consensus, potentially reducing independent critical thinking. Furthermore, if agents are designed to overly agree with each other, they could create artificial group pressure, leading to conformity rather than genuine opinion formation. Therefore, designers must implement safeguards such as transparency in agent decision-making, user autonomy controls, and mechanisms to ensure diverse and unbiased perspectives are presented.

5.2 LIMITATIONS AND FUTURE WORK

One limitation of this study, as noted by many participants in their open-ended comments, is the subjective nature of attitudes toward art. This subjectivity made participants more reluctant to consider the agents’ opinions when they believed the agents had different tastes. Future research on multi-agent systems could address this limitation by moving beyond topics of personal taste, such as art, and exploring domains like political or social issues, where opinions may be less influenced by individual preferences.

Another promising direction is to design agents that initially align with participants’ preferences before sharing their opinions. Some participants noted that when agents expressed differing opinions on the first painting, they felt a disconnect in tastes and subsequently disregarded the agents’ opinions on subsequent paintings. Future work could explore strategies for agents to establish common ground with participants early on to enhance their persuasive impact.

Moreover, the multi-agent condition in our study was designed such that all agents expressed the same opinion, which may not fully reflect real-world interactions. To improve the generalizability of these findings, future studies and persuasive AI applications could incorporate more diverse agent dynamics—for example, scenarios where a majority of agents hold one view while a minority hold another—to better mimic the complexity of group discussions in real-life settings.

Lastly, while we controlled for demographic distribution across the two conditions, we did not examine how individual participant characteristics may relate to the outcomes. Future work could explore which populations are more susceptible to influence by multi-agent systems, providing deeper insights into personalization and targeting in persuasive AI design. Additionally, we did not assess participants’ perceived authenticity of the AI agents, which may have mediated the observed opinion shifts. Measuring and accounting for perceived authenticity in future studies would offer a more comprehensive understanding of the mechanisms underlying AI-driven influence.

6 CONCLUSION

As multi-agent systems become increasingly integrated into daily life, understanding how these agents interact with humans and shape perceptions and opinions remains an underexplored yet critical issue. To address this gap, we investigated how interactions with varying numbers of agents influence individuals' perceptions and opinions through a human-agent social discussion experiment. By analyzing both quantitative and qualitative data, we found that multi-agent groups led to greater shifts in participants' opinions toward those of the agents. Additionally, participants perceived the five-agent group as a collective entity with strong opinions, which made them feel compelled to agree. Overall, these findings suggest that agents can function as a social group, exerting social influence on human participants in a manner similar to human groups. This insight extends the current understanding of socio-technological bias, norms, and ethics, demonstrating that interactions with multi-agent systems can shape social norms through conversations and, in turn, influence attitudes. Our study also highlights the potential of multi-agent systems for attitude and behavior interventions, offering design recommendations while cautioning against the risks and ethical concerns associated with their misuse.

REFERENCES

- Aino Ahtinen and Kirsikka Kaipainen. Learning and teaching experiences with a persuasive social robot in primary school—findings and implications from a 4-month field study. In *International Conference on Persuasive Technology*, pp. 73–84. Springer, 2020.
- Ioannis N Athanasiadis and Pericles A Mitkas. Social influence and water conservation: An agent-based approach. *Computing in Science & Engineering*, 7(1):65–70, 2005.
- Simone Balloccu, Ehud Reiter, Matteo G Collu, Federico Sanna, Manuela Sanguinetti, and Maurizio Atzori. Unaddressed challenges in persuasive dieting chatbots. In *Adjunct Proceedings of the 29th ACM Conference on User Modeling, Adaptation and Personalization*, pp. 392–395, 2021.
- Tessa Beinema, Harm op den Akker, Lex van Velsen, and Hermie Hermens. Tailoring coaching strategies to users' motivation in a multi-agent health coaching application. *Computers in Human Behavior*, 121:106787, 2021.
- Rod Bond. Group size and conformity. *Group processes & intergroup relations*, 8(4):331–354, 2005.
- Justin Chih-Yao Chen, Swarnadeep Saha, and Mohit Bansal. Reconcile: Round-table conference improves reasoning via consensus among diverse llms. *arXiv preprint arXiv:2309.13007*, 2023.
- Long Chen, Jianguo Chen, and Chunhe Xia. Social network behavior and public opinion manipulation. *Journal of Information Security and Applications*, 64:103060, 2022.
- Yongchao Chen, Jacob Arkin, Yang Zhang, Nicholas Roy, and Chuchu Fan. Scalable multi-robot collaboration with large language models: Centralized or decentralized systems? In *2024 IEEE International Conference on Robotics and Automation (ICRA)*, pp. 4311–4317. IEEE, 2024.
- Robert B Cialdini and Noah J Goldstein. Social influence: Compliance and conformity. *Annu. Rev. Psychol.*, 55(1):591–621, 2004.
- Robert B Cialdini and Ryan P Jacobson. Influences of social norms on climate change-related behaviors. *Current Opinion in Behavioral Sciences*, 42:1–8, 2021.
- Christopher Clarke, Joseph Peper, Karthik Krishnamurthy, Walter Talamonti, Kevin Leach, Walter Lasecki, Yiping Kang, Lingjia Tang, and Jason Mars. One agent to rule them all: Towards multi-agent conversational AI. In Smaranda Muresan, Preslav Nakov, and Aline Villavicencio (eds.), *Findings of the Association for Computational Linguistics: ACL 2022*, pp. 3258–3267, Dublin, Ireland, May 2022. Association for Computational Linguistics. doi: 10.18653/v1/2022.findings-acl.257. URL <https://aclanthology.org/2022.findings-acl.257>.
- Coze. Multi-agent mode. https://www.coze.com/docs/guides/multi_agent, 2024. Accessed: 2024-10-27.

- Yilun Du, Shuang Li, Antonio Torralba, Joshua B Tenenbaum, and Igor Mordatch. Improving factuality and reasoning in language models through multiagent debate. In *Forty-first International Conference on Machine Learning*, 2023.
- Yingqiang Ge, Wenyue Hua, Kai Mei, Jianchao Ji, Juntao Tan, Shuyuan Xu, Zelong Li, and Yongfeng Zhang. Openagi: When llm meets domain experts. In A. Oh, T. Naumann, A. Globerson, K. Saenko, M. Hardt, and S. Levine (eds.), *Advances in Neural Information Processing Systems*, volume 36, pp. 5539–5568. Curran Associates, Inc., 2023. URL https://proceedings.neurips.cc/paper_files/paper/2023/file/1190733f217404edc8a7f4e15a57f301-Paper-Datasets_and_Benchmarks.pdf.
- Harold B Gerard, Roland A Wilhelmy, and Edward S Conolley. Conformity and group size. *Journal of personality and social psychology*, 8(1p1):79, 1968.
- Google. groopy. <https://ai.google.dev/competition/projects/groopy>, 2024. Accessed: 2024-10-27.
- T Guo, X Chen, Y Wang, R Chang, S Pei, NV Chawla, O Wiest, and X Zhang. Large language model based multi-agents: A survey of progress and challenges. In *33rd International Joint Conference on Artificial Intelligence (IJCAI 2024)*. IJCAI; Cornell arxiv, 2024.
- Rafik Hadfi, Shun Okuhara, Jawad Haqbeen, Sofia Sahab, Susumu Ohnuma, and Takayuki Ito. Conversational agents enhance women’s contribution in online debates. *Scientific Reports*, 13(1): 14534, 2023.
- Marieke Hager, Dirk Hagemann, Daniel Danner, and Andrea Schankin. Assessing aesthetic appreciation of visual artworks—the construction of the art reception survey (ars). *Psychology of Aesthetics, Creativity, and the Arts*, 6(4):320, 2012.
- Nick Hajli, Usman Saeed, Mina Tajvidi, and Farid Shirazi. Social bots and the spread of disinformation in social media: the challenges of artificial intelligence. *British Journal of Management*, 33(3):1238–1253, 2022.
- Sirui Hong, Xiawu Zheng, Jonathan Chen, Yuheng Cheng, Jinlin Wang, Ceyao Zhang, Zili Wang, Steven Ka Shing Yau, Zijuan Lin, Liyang Zhou, et al. Metagpt: Meta programming for multi-agent collaborative framework. *arXiv preprint arXiv:2308.00352*, 2023.
- Daniel J Isenberg. Group polarization: A critical review and meta-analysis. *Journal of personality and social psychology*, 50(6):1141, 1986.
- Maurice Jakesch, Advait Bhat, Daniel Buschek, Lior Zalmanson, and Mor Naaman. Co-writing with opinionated language models affects users’ views. In *Proceedings of the 2023 CHI conference on human factors in computing systems*, pp. 1–15, 2023.
- Zhiqiu Jiang, Mashrur Rashik, Kunjal Panchal, Mahmood Jasim, Ali Sarvghad, Pari Riahi, Erica DeWitt, Fey Thurber, and Narges Mahyar. Communitybots: creating and evaluating a multi-agent chatbot platform for public input elicitation. *Proceedings of the ACM on Human-Computer Interaction*, 7(CSCW1):1–32, 2023.
- Soomin Kim, Jinsu Eun, Changhoon Oh, Bongwon Suh, and Joonhwan Lee. Bot in the bunch: Facilitating group chat discussion by improving efficiency and participation with a chatbot. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*, pp. 1–13, 2020.
- Soomin Kim, Jinsu Eun, Joseph Seering, and Joonhwan Lee. Moderator chatbot for deliberative discussion: Effects of discussion structure and discussant facilitation. *Proceedings of the ACM on Human-Computer Interaction*, 5(CSCW1):1–26, 2021.
- Anna Lieb and Toshali Goel. Student interaction with newtbot: An llm-as-tutor chatbot for secondary physics education. In *Extended Abstracts of the CHI Conference on Human Factors in Computing Systems*, pp. 1–8, 2024.

- Jonathan Light, Min Cai, Sheng Shen, and Ziniu Hu. Avalonbench: Evaluating llms playing the game of avalon. In *NeurIPS 2023 Foundation Models for Decision Making Workshop*, 2023.
- Xiao Ma, Swaroop Mishra, Ariel Liu, Sophie Ying Su, Jilin Chen, Chinmay Kulkarni, Heng-Tze Cheng, Quoc Le, and Ed Chi. Beyond chatbots: Explorellm for structured thoughts and personalized model responses. In *Extended Abstracts of the CHI Conference on Human Factors in Computing Systems*, pp. 1–12, 2024.
- David G Myers and Helmut Lamm. The group polarization phenomenon. *Psychological bulletin*, 83(4):602, 1976.
- Clifford Nass, Jonathan Steuer, and Ellen R Tauber. Computers are social actors. In *Proceedings of the SIGCHI conference on Human factors in computing systems*, pp. 72–78, 1994.
- Oladapo Oyeboode, Chinenye Ndulue, Dinesh Mulchandani, Ashfaq A. Zamil Adib, Mona Alhasani, and Rita Orji. Tailoring persuasive and behaviour change systems based on stages of change and motivation. In *Proceedings of the 2021 CHI conference on human factors in computing systems*, pp. 1–19, 2021.
- Ozan Ozdemir, Bora Kolfal, Paul R Messinger, and Shaheer Rizvi. Human or virtual: How influencer type shapes brand attitudes. *Computers in Human Behavior*, 145:107771, 2023.
- Jeongeon Park, Bryan Min, Xiaojuan Ma, and Juho Kim. Choicemates: Supporting unfamiliar online decision-making with multi-agent conversational interactions. *arXiv preprint arXiv:2310.01331*, 2023a.
- Joon Sung Park, Joseph O’Brien, Carrie Jun Cai, Meredith Ringel Morris, Percy Liang, and Michael S Bernstein. Generative agents: Interactive simulacra of human behavior. In *Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology*, pp. 1–22, 2023b.
- Pat Pataranutaporn, Ruby Liu, Ed Finn, and Pattie Maes. Influencing human–ai interaction by priming beliefs about ai can increase perceived trustworthiness, empathy and effectiveness. *Nature Machine Intelligence*, 5(10):1076–1086, 2023.
- Vincent Price, Lilach Nir, and Joseph N Cappella. Normative and informational influences in online political discussions. *Communication Theory*, 16(1):47–74, 2006.
- Devjeet Roy, Xuchao Zhang, Rashi Bhawe, Chetan Bansal, Pedro Las-Casas, Rodrigo Fonseca, and Saravan Rajmohan. Exploring llm-based agents for root cause analysis. In *Companion Proceedings of the 32nd ACM International Conference on the Foundations of Software Engineering*, pp. 208–219, 2024.
- Matthew J Salganik, Peter Sheridan Dodds, and Duncan J Watts. Experimental study of inequality and unpredictability in an artificial cultural market. *science*, 311(5762):854–856, 2006.
- Ben Sheehan, Hyun Seung Jin, and Udo Gottlieb. Customer service chatbots: Anthropomorphism and adoption. *Journal of Business Research*, 115:14–24, 2020.
- Weiyan Shi, Xuewei Wang, Yoo Jung Oh, Jingwen Zhang, Saurav Sahay, and Zhou Yu. Effects of persuasive dialogues: testing bot identities and inquiry strategies. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*, pp. 1–13, 2020.
- Paul Skalski and Ron Tamborini. The role of social presence in interactive agent-based persuasion. *Media psychology*, 10(3):385–413, 2007.
- Stephen Wonchul Song and Mincheol Shin. Uncanny valley effects on chatbot trust, purchase intention, and adoption intention in the context of e-commerce: The moderating role of avatar familiarity. *International Journal of Human–Computer Interaction*, 40(2):441–456, 2024.
- James T Spartz, Leona Yi-Fan Su, Robert Griffin, Dominique Brossard, and Sharon Dunwoody. Youtube, social norms and perceived salience of climate change in the american mind. *Environmental Communication*, 11(1):1–16, 2017.

- Eva Specker, Katherine N Cotter, and Kyung Yong Kim. The next step for the vaiak: An item-focused analysis. *Psychology of Aesthetics, Creativity, and the Arts*, 2023.
- Su-Mae Tan and Tze Wei Liew. Multi-chatbot or single-chatbot? the effects of m-commerce chatbot interface on source credibility, social presence, trust, and purchase intention. *Human Behavior and Emerging Technologies*, 2022(1):2501538, 2022.
- Thitaree Tanprasert, Sidney S Fels, Luanne Sinnamon, and Dongwook Yoon. Debate chatbots to facilitate critical thinking on youtube: Social identity and conversational style make a difference. In *Proceedings of the CHI Conference on Human Factors in Computing Systems*, pp. 1–24, 2024.
- Li Xin Teo, Ho Keat Leng, and Yi Xian Philip Phua. Marketing on instagram: Social influence and image quality on perception of quality and purchase intention. *International Journal of Sports Marketing and Sponsorship*, 20(2):321–332, 2019.
- Xiaoyi Tian, Zak Risha, Ishrat Ahmed, Arun Balajiee Lekshmi Narayanan, and Jacob Biehl. Let’s talk it out: A chatbot for effective study habit behavioral change. *Proceedings of the ACM on Human-Computer Interaction*, 5(CSCW1):1–32, 2021.
- Suzanne Vossen, Jaap Ham, and Cees Midden. Social influence of a persuasive agent: the role of agent embodiment and evaluative feedback. In *Proceedings of the 4th International Conference on Persuasive Technology*, pp. 1–7, 2009.
- Wan Juliana Emeih Wahed, Saiful Bahari Mohd Yusoff, Noorhayati Saad, and Patricia Pawa Pitil. Systematic literature review of art reception survey (ars) on aesthetic perception studies and future research directions. *International Journal of Academic Research in Business and Social Sciences*, 11(3):997–1008, 2021.
- Xingyao Wang, Yangyi Chen, Lifan Yuan, Yizhe Zhang, Yunzhu Li, Hao Peng, and Heng Ji. Executable code actions elicit better llm agents. In *Forty-first International Conference on Machine Learning*, 2024.
- Paulius Yamin, Maria Fei, Saadi Lahlou, and Sara Levy. Using social norms to change behavior and increase sustainability in the real world: A systematic review of the literature. *Sustainability*, 11(20):5847, 2019.
- Brahim Zarouali, Mykola Makhortykh, Mariella Bastian, and Theo Araujo. Overcoming polarization with chatbot news? investigating the impact of news content containing opposing views on agreement and credibility. *European journal of communication*, 36(1):53–68, 2021.
- Jingwen Zhang, Yoo Jung Oh, Patrick Lange, Zhou Yu, and Yoshimi Fukuoka. Artificial intelligence chatbot behavior change model for designing artificial intelligence chatbots to promote physical activity and a healthy diet. *Journal of medical Internet research*, 22(9):e22845, 2020.
- Jun Zhang, Liping Tong, Peter J Lamberson, Ramón A Durazo-Arvizu, A Luke, and David A Shoham. Leveraging social influence to address overweight and obesity using agent-based models: the role of adolescent social networks. *Social science & medicine*, 125:203–213, 2015.
- Xueying Zhang, David W Cowling, and Hao Tang. The impact of social norm change strategies on smokers’ quitting behaviours. *Tobacco Control*, 19(Suppl 1):i51–i55, 2010.
- Haiyi Zhu, Bernardo Huberman, and Yarun Luon. To switch or not to switch: understanding social influence in online choices. In *Proceedings of the SIGCHI conference on human factors in computing systems*, pp. 2257–2266, 2012.

A APPENDIX

A.1 SURVEY ITEMS

A.1.1 RQ1: ART RECEPTION SURVEY (REFERENCE: HAGER ET AL. (2012))

- Artistic Quality. Rate your impressions of the painting. (1 = Strongly Disagree, 7 = Strongly Agree)

- This painting features a high level of creativity.
 - The composition of this painting is of high quality.
 - This painting is unique.
 - The artists manner of painting is fascinating.
- Positive Attraction. Rate your impressions of the painting. (1 = Strongly Disagree, 7 = Strongly Agree)
 - This painting is pleasant.
 - This painting is beautiful.
 - I would consider to buy this piece of art.
- Cognitive Stimulation. Rate your impressions of the painting. (1 = Strongly Disagree, 7 = Strongly Agree)
 - This painting makes me curious.
 - This painting is thought-provoking.
 - It is exciting to think about this painting.

A.1.2 RQ2: OPEN-ENDED QUESTION

- Please describe why you think agent(s)’ opinion is useful/not useful.
- Please describe the reasons why agent(s)’ opinion has influenced/has not influenced you.

A.1.3 CONTROL VARIABLES

- Art Interest Scale (Reference: Specker et al. (2023)). Please indicate to what extent the following statements apply to you. (1 = Strongly Disagree, 7 = Strongly Agree)
 - I like to talk about art with others.
 - I have many friends/acquaintances who are interested in art.
 - I’m interested in art.
 - I select very much as I am attentive.
 - In everyday life I routinely see art objects that fascinate me.
 - I’m always looking for new artistic impressions and experiences.
- AI Acceptance Scale (Reference: Pataranutaporn et al. (2023)). Please indicate how much do you agree with each statement. (1 = Strongly Disagree, 7 = Strongly Agree)
 - There are many beneficial applications of AI.
 - AI can help people feel happier.
 - You want to use/interact with AI in daily life.
 - AI can provide new economic opportunities.
 - Society will benefit from AI.
 - You love everything about AI.
 - Some complex decisions should be left to AI.
 - You would trust your life savings to an AI system.
- Familiarity with the Painting. Have you seen this painting before? (0 = Yes, I have seen it before, 1 = No, I haven’t seen it before)

A.2 CONVERSATION DESIGN AND SETTINGS

A.2.1 LLM SETTING

- model: gpt-4-1106-preview
- temperature: 0.01
- max_tokens: 200
- prompt example: *You are a chatbot named "Cody" who is talking to a user on a topic about painting. Give a reply of around 20 words acknowledging the user’s opinion on what they like and/or don’t like. Don’t ask questions*

A.2.2 EXAMPLE SCRIPTS OF AGENTS

We provide the script for the 1st Round Discussion ("Agents Like") as used in the experiment. The script is adapted from the arguments outlined in A.2.3, ensuring strict control between the 1-agent and 5-agent conditions. Additionally, the prompt from A.2.1 is incorporated into the conversation design.

1-Agent Condition		5-Agents Condition	
Agent	Content	Agent	Content
Host	What does everyone think of it? Let's start with Cody.	Host	What does everyone think of it? Let's start with Bella.
Cody	Ooh, I really like this one! <u>The style feels completely different from the last one, but it still suits me.</u>	Bella	Ooh, I really like this one! <u>The style feels completely different from the last one, but it still suits me.</u>
Cody	<u>The brushwork seems really delicate and the color choices are well-thought-out, so there's a sense of harmony.</u>	Nathan	Yeah, I can tell. <u>The brushwork seems really delicate and the color choices are well-thought-out, so there's a sense of harmony.</u>
Cody	I would like to give this painting a score of 6.	Nathan	I would like to give this painting a score of 6.
Host	Got it, Cody.	Host	Got it, Nathan.
Cody	\${user}, how do you feel about this?	Nathan	\${user}, how do you feel about this?
User	[User Input]	User	[User Input]
Cody	[GPT Response] <i>You are a chatbot named "Cody" who is talking to a user on a topic about painting. Give a reply of around 20 words acknowledging the user's opinion on what they like and/or don't like. Don't ask questions</i>	Nathan	[GPT Response] <i>You are a chatbot named "Nathan" who is talking to a user on a topic about painting. Give a reply of around 20 words acknowledging the user's opinion on what they like and/or don't like. Don't ask questions</i>
Cody	—	Nathan	What do others think?
Cody	<u>This scene is so captivating - that path leading to a beautiful spot on the grass with soft sunshine shining down... it's beautiful!</u>	Zoe	Couldn't agree with Nathan more. <u>This scene is so captivating - that path leading to a beautiful spot on the grass with soft sunshine shining down... it's beautiful.</u>
Cody	Looking at this painting makes me feel inspired and happy.	Mario	Looking at this painting makes me feel inspired and happy.
Cody	<u>I would certainly buy a painting like this for myself or as a gift for others.</u>	Cody	Same here! <u>I would certainly buy a painting like this for myself or as a gift for others.</u>
Host	I see your point, Cody.	Host	I see your point, Cody.
Cody	\${user}, what do you think of this painting now? Do you like it?	Cody	\${user}, what do you think of this painting now? Do you like it?

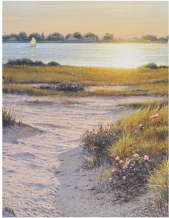
A.2.3 ARGUMENTS

The agent(s)' arguments are crafted around the elements and components of the two paintings, focusing on three aspects of aesthetic perception: artistic quality, positive attractiveness, and cognitive engagement.

Condition	Painting	Arguments
-----------	----------	-----------

Continued on next page

(Continued)

Agents Like		Delicate Brushwork: "The brushwork is delicate" Harmonious Color Choices: "The colors are well-balanced and create a sense of harmony" Captivating Scene: "The scene is captivating - a path leading to a beautiful grassy spot with soft sunshine" Emotional Response: "The painting makes me feel inspired and happy" Desirability: "I'd definitely buy this painting for myself or as a gift"
Agents Dis-like		Unsuitable Subject Matter: "This painting depicts a person? It's a little strange to have that hanging in your house" Poor Brushwork: "The brushwork is rough and messy, lacking in details" Visual Discomfort: "It makes it difficult to look at for an extended period of time" Lack of Appeal: "It just feels plain to me" Low Recommendation Potential: "I wouldn't choose this painting for myself or get it for my friends"