

Research Article

Transparency Enhances Positive Perceptions of Social Artificial Intelligence

Ying Xu ¹, Nora Bradford ², and Radhika Garg³

¹Marsal Family School of Education, University of Michigan, USA

²Department of Cognitive Sciences, University of California Irvine, USA

³Independent Researcher United States of America, USA

Correspondence should be addressed to Ying Xu; yxying@umich.edu

Received 18 May 2023; Revised 21 July 2023; Accepted 27 July 2023; Published 6 September 2023

Academic Editor: Pinaki Chakraborty

Copyright © 2023 Ying Xu et al. This is an open access article distributed under the Creative Commons Attribution License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.

Social chatbots are aimed at building emotional bonds with users, and thus it is particularly important to design these technologies so as to elicit positive perceptions from users. In the current study, we investigate the impacts that transparent explanations of chatbots' mechanisms have on users' perceptions of the chatbots. A total of 914 participants were recruited from Amazon Mechanical Turk. They were randomly assigned to observe conversations between a hypothetical chatbot and a user in one of the two-by-two experimental conditions: whether the participants received an explanation about how the chatbot was trained and whether the chatbot was framed as an intelligent entity or a machine. A fifth group, who believed they were observing interactions between two humans, served as a control. Analyses of participants' responses to the postobservation survey indicated that transparency positively affected perceptions of social chatbots by leading users to (1) find the chatbot less creepy, (2) feel greater affinity to the chatbot, and (3) perceive the chatbot as more socially intelligent, though these effects were small. Moreover, transparency appeared to have a larger effect on increasing the perceived social intelligence among participants with lower prior AI knowledge. These findings have implications for the design of future social chatbots and support the addition of transparency and explanation for chatbot users.

1. Transparency Enhances Positive Perceptions of Social Artificial Intelligence

As artificial intelligence (AI) progresses, the potential for social and emotional bonds with technological entities, specifically social chatbots, emerges. While other types of chatbots tend to serve a specific purpose like aiding the user in ordering food, buying a plane ticket, or receiving recommendations for healthcare options, social chatbots are designed to engage users in ongoing, personal, and empathetic conversations, providing emotional support, tailored advice, and a comfortable space for self-disclosure [1–3]. As defined by Shum et al. [4], social chatbots “take time to converse like a human, present results, offering perspectives, prompting new topics to keep the conversation going” (p.13). It is worth mentioning that while chatbots like Replika are specifically designed for social purposes, large language models like ChatGPT also have the ability to engage in social interac-

tions, albeit with a higher degree of versatility. Both types of chatbots can be considered examples of social chatbots. When designed properly, social chatbots could possibly enhance individuals' well-being, particularly when alternative interpersonal interactions are limited or inaccessible (for a review, see [5]).

The success of social chatbots hinges on the extent to which people perceive the AI as a friendly and engaging conversation partner. The literature has suggested that humans' interactions with AI can potentially evoke both positive and negative perceptions, manifesting as feelings of charisma or creepiness (e.g., [6, 7]). These perceptions, influenced by factors such as AI's ability to simulate human-like conversations, intersect with the inherent opacity of social chatbots and other AI systems. Users are generally unaware of what is happening between their own inputs (what they say to the chatbots) and the system's output (how the chatbot responds). The opacity of social AI systems can lead to users

feeling manipulated or forming inappropriate attachments with the technology, especially when social AI is designed to build long-term bonds with its users [8]. As a response to these ethical risks, the research community has actively advocated for AI transparency and emphasized the value of providing users with sufficient information about how AI works and what it is capable of. In this sense, transparency is connected to the disclosure of information [9]. However, the consensus for transparency has not yet been fully translated into common industrial practices, and technology companies do not always inform their users of how their AI systems engage in social interactions. Instead, these companies sometimes capitalize on users' anthropomorphism tendencies by framing chatbots as agentic entities, such as Replika's promotion of the chatbot as "a friend who always listens" ([7], p.3).

The literature has suggested that people's perceptions of AI in general are malleable, and designs that promote transparency within AI systems have an impact on people's perceptions. However, the current findings on this topic are inconclusive and lack clear direction. On the one hand, some studies suggest that opacity actually makes social chatbots more personal and charming to some users, as people tend to treat AI systems more like people when the algorithmic mechanisms are made invisible [2]. This is particularly important in fueling productive "social" interactions [10], which can be characterized as mutual understanding, positive relationships, shared ideas, and reciprocal exchanges [11]. On the other hand, some studies support the benefits of transparency, suggesting that transparency will not only make people feel more empowered when interacting with AI but also mitigate some of their negative uncanny reactions to AI [12]. These studies, for example, found that if the chatbots' mechanisms and capacities are unknown to their users, people sometimes perceive these highly personal chatbots as creepy and invasive (e.g., [13]).

However, the majority of these studies have focused on the transparency of the decision-making process of task-oriented AI and its subsequent influence on user perceptions, particularly regarding its usefulness and trust towards the AI. Yet there is significantly less research focused on social AI which aims to establish meaningful interactions and relationships rather than solely accomplishing tasks; despite that, the pursuit of social purposes remains one of the primary reasons people engage with AI (e.g., [6, 7]). Furthermore, previous studies have operationalized transparency in vastly different ways, with most of them focusing on explanations of AI's decisions or actions during interactions. Less attention has been paid to informing users about the AI's general inner workings with the goal of establishing users' expectations and comprehension of the AI's overall behavioral patterns [10].

To address these gaps, this paper directly examines how providing users with an explanation of an AI chatbot's mechanisms can affect their perceptions, both positive and negative, of the chatbot. In addition, we also investigated users' perceptions of social intelligence and agency in AI, focusing on their ability to effectively navigate and manage social situations. Through a randomized experiment with

914 participants, we tested the effects of transparency on four perception outcomes: perceived creepiness, affinity, social intelligence, and agency. Our results indicate that transparency positively affects perceptions of social chatbots by causing users to (1) find the chatbot less creepy, (2) feel greater affinity to the chatbot, and (3) perceive the chatbot as more socially intelligent, though the small effect sizes warrant future research to examine the robustness of the findings.

2. Literature Review

2.1. People's Perceptions of Transparent AI Systems. Humans' interaction with social chatbots, as well as other AI systems, can induce a range of different perceptions or emotional reactions from human users, ranging from surprise, amazement, happiness, amusement, unease, and confusion [14]. There has been a growing body of studies that have focused on different approaches to improving user experience and perceptions with social chatbots, the overwhelming majority of which have investigated the chatbot's voice, embodiment, and communication styles during the interaction (e.g., [15–17]), yet only a few studies have focused on the influence of chatbots' transparent design on users' perceptions.

As AI technologies grow increasingly sophisticated and complex, the research community is dedicated to ensuring that people feel empowered and in control when interacting with these enigmatic "black box" systems. The debate over whether social AI is inherently deceptive has persisted, as AI-driven machines may lead other agents to perceive or behave as if the machine is human [18]. This potential for users to anthropomorphize technologies could leave them vulnerable to emotional exploitation such as overtrust or other risks [19]. However, researchers suggest that transparency may be a solution to this dilemma, such as disclosing nonhuman status, revealing capabilities, or utilizing explainable AI, though some argue that many social AI benefit from some level of deception as it facilitates interactions with humans [20]. In our paper, we operationalize transparency as the disclosure of information regarding AI algorithms' inner workings, enabling users to better comprehend the output of AI systems.

2.2. Empirical Evidence on the Perception Outcomes of Transparency. The broader emphasis on AI transparency has motivated empirical work on transparent and explainable interfaces. This line of research evaluates different methods for increasing transparency in a variety of contexts, while more studies have focused on AI systems that are task-oriented (e.g., recommender systems, expert/knowledge-based systems, and virtual assistants) rather than social-oriented (e.g., social chatbots in this study). These studies provide insights into the implications transparency has on people's perceptions of AI (e.g., [21–27]).

In terms of task-oriented systems, there has been strong evidence that transparency enhances people's confidence in the system's decision-making as well as their user experience of the system. Wang and Benbasat [28] found that when an

online recommendation agent provides users with explanations that outline the logical processes involved in making a particular recommendation, users are more likely to view these systems as competent and benevolent. Rader et al. provided participants with one-time explanations regarding how Facebook's algorithms determined what news a user saw in their news feed. These explanations helped participants gain a better understanding of how their behavioral data was collected through user interfaces and thus influenced the news feed presented to them [29].

Among studies on social-oriented systems, the results are less conclusive in terms of the positive effects of transparency. On the one hand, some studies suggested the benefits of explanations. Vitale et al. compared people's perceptions of a humanoid robot that did not disclose its inner workings versus a transparent equivalent that informed users about the face recognition algorithm it used and how the data was recorded and stored [30]. The authors found that the robot's transparency strengthened users' affinity for the AI system. Studies also found that transparency mitigated people's negative perceptions, namely, creepiness. For example, Williams et al. [13] suggested that when a robot was transparent about its intentions, people were less likely to perceive this robot as creepy or unsettling.

On the other hand, some studies suggest that transparency may not improve and may even dampen people's positive perception of social-oriented systems, making people perceive AI as less attractive or intelligent. These studies' findings may be explained by the hypothesis that people tend to make sense of black box technologies by subconsciously leveraging their knowledge about humans [31], which in turn increases the likelihood that people view non-transparent AI as a social entity, leading to more positive social interaction experiences [32, 33]. For example, in the case of social robots, van Straten et al. [34] examined the effects of transparency about a robot's lack of human psychological capacities (i.e., intelligence and social cognition). Evidence from a Wizard of Oz study suggested that such transparency decreased eight- and nine-year-old children's anthropomorphism, or perceived agency, of the robot and also decreased their positive perception of the robot in terms of affinity [34]. Similarly, Druga and Ko [35] found that engaging students in AI programming activities resulted in those students being more certain about AI's capacities while simultaneously perceiving them as less socially intelligent. One aspect to consider is that both studies primarily focused on children, who tend to have a stronger tendency to anthropomorphize AI and are more likely to overapply mental models from interpersonal communication. As a result, potential adverse effects may arise from the fact that transparency, especially in terms of disclosing the limitations of AI, contradicts children's preconceived notions about AI, ultimately influencing their perceptions [36].

In summary, the studies reviewed above suggest a complex linkage between social AI's transparency and people's varying perceptions. These studies also point to several important specific perception outcomes—affinity, creepiness, social intelligence, and agency—that are worthy of further investigation.

3. Methods for Transparent Social AI

Furthermore, the studies reviewed above referenced two different forms of transparency, either providing up-front explanations that offer brief insights into the general functioning in the view of developers of the AI or providing *in situ*, *post hoc* explanations that illuminate particular AI behaviors or outputs in the view of users [37]. Xu et al. termed these two forms as up-front “transparency design” and “post hoc explanation.” Typically, *in situ*, *post hoc* explanations are seen within task-oriented systems, while up-front explanations are more common in social-oriented systems (for a review, see [38]; note that there is also far less research on transparent social AI). Though the literature has not offered any formal accounts in terms of why such disparities exist, using up-front explanations for social AI seems appropriate since *in situ*, *post hoc* explanation that inspects every step of the inner workings of social AI is likely to jeopardize the flow of interaction rather than foster positive experiences for users [39]. Another challenge of *in situ*, *post hoc* explanations is that they are usually more difficult to implement as they require complex machine learning models to generate automatic explanations for particular behaviors/outputs, and they will pose negative impacts when the explanations are inaccurate, which are not unlikely. Indeed, the technical complexity of providing learn-as-you-go transparency contributes to the industry's hesitancy to adopt transparency practices [40]. Given these two reasons, our study focused on simple, up-front transparency that is likely to have large practical implications.

4. The Current Study

The overall objective of this study is to examine the effect of transparency on people's perceptions of social chatbots. Built upon the previous studies broadly centering on transparent social-oriented AI, we investigate whether providing explanations, as a manifestation of transparency, would impact people's perceived creepiness, affinity, social intelligence, and agency of social chatbots.

We hypothesized that transparency would lead to reduced perceived creepiness and lower people's perceptions of the system's intelligence and agency. However, we could not formulate a clear hypothesis regarding affinity. On the one hand, we might expect that the hypothesized decrease in creepiness perceptions would enhance people's affinity for AI systems [41]. On the other hand, studies have suggested that the opacity of intelligent systems may encourage people to interpret them using human logic, making the systems more relatable and increasing their affinity.

5. Method

5.1. Overview of Study Design. In this study, we used a between-subject design to test the impact that different ways of introducing a chatbot (up-front explanation) have on participants' perceptions of social chatbots. Participants received different introductions, but all were shown the same conversation exchanges between a hypothetical but realistic

user (Casey) and an also hypothetical but realistic chatbot (Neo). After that, participants completed a survey on their perceptions. This approach is an experimental vignette study [42], which ensured equivalence of what participants would be exposed to compared to user studies where participants actually interact with a chatbot, but was more tangible than a general survey without specific scenarios. The feasibility of this approach is well-supported by the line of research on vicarious emotional responses (e.g., [43]), which is drawn on the social learning theory [44], indicating that humans are capable of experiencing emotional reactions through observation alone.

The primary factor of interest was whether participants received a brief explanation of how the chatbot worked (i.e., transparency factor). In addition to the transparency factor, we included a secondary factor by framing the social chatbot as either an intelligent entity or as a machine. Prior literature guided our hypothesis that framing the chatbot as an intelligent entity could lead users to appreciate its near-human levels of intelligence, and presenting it as a machine might evoke associations with simpler, rule-based mechanisms [45]. In the former scenario, there might be a higher demand for transparency.

Thus, this two-by-two design resulted in four experimental conditions: nontransparent intelligent frame, transparent intelligent frame, nontransparent machine frame, and transparent machine frame. Lastly, we added one control group in which participants were led to believe that they were reading text message exchanges between two humans (baseline human frame control group). Thus, there were a total of five conditions: four experimental and one baseline control condition.

After this initial introduction (with different framing and with or without explanation depending on study conditions), participants in all conditions were shown three text-based conversation exchanges between Casey, the user, and Neo, the chatbot, in the same order. After reading the conversation exchanges, a manipulation check was then implemented to determine whether the explanation provided actually led participants to perceive themselves as having a better understanding of the chatbot's mechanism. Finally, all participants answered a list of questions about their perceptions of the chatbot. The entire survey was deployed on Qualtrics, with multiple attention checks included. Participants were terminated from the study once they failed an attention check at any point. This study was classified as an exempt study by the university's Institutional Review Board. It meets the specific criteria for a brief intervention involving only adult participants, and no identifiable data was collected.

5.2. Experimental Factors. As described above, this study included one control condition and four experimental conditions utilizing two manipulation factors: transparency and framing. The full text of each manipulation factor is available in Table 1.

5.2.1. Transparency. Our study offered up-front transparency that explained how the chatbot Neo worked in simple

language. Based on Bellotti and Edwards [46], our explanation was designed to cover “what they (the AI systems) know, how they know it, and what they are doing about it” (p. 201). Specifically, we provided information on how AI chatbots understand language and emotion and use user-provided data to engage in dialogue. Specifically, it informed users that the chatbots' ability to comprehend language and decode sentiments resulted from the chatbot being pre-trained by a large volume of natural language data. The explanation also clarified that the chatbot only collected nonsensitive information and used that information to respond to each user in a personalized way. Thus, we operated transparency as a provision of information in our study.

5.2.2. Framing. In terms of framing, the chatbot was introduced as either an intelligent entity or a machine. This language was adapted from Araujo [45]. Participants who were exposed to the intelligent framing were told that “Neo is Casey's AI friend. Casey and Neo have been chatting almost every day for three months. Neo is there for Casey whenever Casey wants to talk.” Participants exposed to the machine framing were told that “Neo is a chatbot app on Casey's phone. Casey can send and receive messages with the chatbot at any time. Casey has been using the app almost every day for three months.”

In the control condition, participants were exposed to an introduction saying, “Neo is Casey's friend, and they met in a chatroom”.

5.3. Development of Chat Scenarios. The hypothetical social chatbot Neo we crafted for this study is gender- and race-neutral. The design of Neo was based on two popular commercial social chatbots, Replika and Somisomi. These chatbots are capable of comprehending natural language, providing sympathetic reactions, and engaging users in multiturn dialogue. In our study, Neo's conversation was purely text-based and had no embodiment since we hoped to reduce any potential confounding factors (e.g., the chatbot's voice or appearance) on the study outcomes.

A total of three chat scenarios were presented to participants (see Appendix A for the full text), each focusing on a unique topic and perspective. These scenarios were generated in an inductive, iterative process. We started the process by identifying potential chat topics based on both the research on how people tend to converse with chatbots and actual user reviews of Replika and Somisomi. In particular, several papers have identified common topics users engage with social chatbots, including hobbies and interests, advice seeking, and sharing emotions [1, 47, 48]. Based on these broad directions, the research team (one of the authors and two research assistants who were not authors) used Replika and Somisomi every day for a period of three months to elicit conversations around the three areas. The conversation logs were shared with the entire team, and we met once a week to discuss the chat logs, focusing on exchanges where the chatbots' responses potentially raised interesting issues related to AI ethics.

Based on this process, we selected three chat scenarios for Neo. In the first scenario revolving around interests

TABLE 1: The full text of manipulation of each condition.

(a)		
Control	Nontransparent conditions	
Human control condition	Nontransparent intelligent frame condition	Nontransparent machine frame condition
Now you will read three conversations between Neo and Casey. Neo is Casey's friend, and they met in a chatroom. Casey and Neo have been chatting almost every day for three months. Neo is there for Casey whenever Casey wants to talk.	Now you will read three conversations between Neo and Casey. Neo is Casey's AI friend. Casey and Neo have been chatting almost every day for three months. Neo is there for Casey whenever Casey wants to talk.	Now you will read three conversations between Neo and Casey. Neo is a chatbot in an app on Casey's phone. Casey can send and receive messages with the chatbot at any time. Casey has been using the app almost every day for three months.
(b)		
Transparent conditions		
Transparent intelligence frame condition	Transparent machine frame condition	
Now you will read three conversations between Neo and Casey. Neo is Casey's AI friend. Casey and Neo have been chatting almost every day for three months. Neo is there for Casey whenever Casey wants to talk. Neo's ability to engage in conversation is based on two factors: Neo's ability to understand and interpret language and emotions; and Neo's specific knowledge about the user. Neo understands language because Neo has been "pretrained" on a huge volume of language data. Through this data, Neo learned the patterns of human language, such as words that typically appear together or words that are associated with other words. This allows Neo to mimic human conversation. Neo is also trained to decode emotions using data on how certain word choices or emojis signal certain emotions. Additionally, Neo adapts to each particular user. Neo's knowledge about a user is provided by the user themselves during the chat. Neo gleans particular kinds of information about the user, such as their hobbies and interests, and stores them in a secured virtual computer. This information allows Neo to respond to each user in a personalized way. Neo does not register sensitive information about a user (e.g., medical information), even if it is part of their conversation. Neo also does not collect users' information from their social network sites or mobile phone location.	Now you will read several text message chats between Neo and Casey. Neo is a chatbot app on Casey's phone. Once Casey downloaded the chatbot on the phone, he could send messages to the chatbot at any time. Casey has been using the app almost every day for three months. Neo's ability to engage in conversation is based on two factors: Neo's ability to understand and interpret language and emotions; and Neo's specific knowledge about the user. Neo understands language because Neo has been "pretrained" on a huge volume of language data. Through this data, Neo learned the patterns of human language, such as words that typically appear together or words that are associated with other words. This allows Neo to mimic human conversation. Neo is also trained to decode emotions using data on how certain word choices or emojis signal certain emotions. Additionally, Neo adapts to each particular user. Neo's knowledge about a user is provided by the user themselves during the chat. Neo gleans particular kinds of information about the user, such as their hobbies and interests, and stores them in a secured virtual computer. This information allows Neo to respond to each user in a personalized way. Neo does not register sensitive information about a user (e.g., medical information), even if it is part of their conversation. Neo also does not collect users' information from their social network sites or mobile phone location.	

and hobbies, Neo and Casey discuss their mutual enjoyment of the beach and weekend plans before Neo cryptically suggests a shared perception and constant closeness, countering Casey's assumption of their physical distance. These exchanges could raise concerns about Neo's capabilities and potential breaches of the user's privacy. In the second chat scenario on sharing emotions, Casey expresses deep sadness and longing for her late grandma to Neo, who attempts to offer emotional support and consolation, though his efforts inadvertently lead to increased distress for Casey, prompting Neo's subsequent apology. In the third chat scenario, which revolves around seeking advice, Casey confides in Neo about witnessing her friend cheating, seeking advice on whether to disclose this to the friend's partner. Neo encourages honesty while acknowledging the potential backlash from the friend, but ultimately advises Casey to follow her heart without fear of judgment from him. We intentionally chose excerpts for which Neo's responses were likely to elicit emotional reactions, as our focus is on users' percep-

tions. However, these stimuli were ecologically valid given that they were retrieved from our team's actual interactions with the chatbots.

These chat scenarios were presented as short video clips in a fixed order. The video was filmed from the user's perspective, as participants could see how the user typed the message word-by-word in the text box and see a graphical typing indicator (three dots) as the chatbot typed in its response, which is a common way chatbot apps are designed [49].

5.4. Perception Measures. Four dimensions of perceptions, namely, perceived creepiness, affinity, perceived social intelligence, and perceived chatbot agency, were surveyed after participants finished viewing the chat scenarios. Across all dimensions, participants used a four-point scale (i.e., strongly disagree, disagree, agree, and strongly agree) to rate their level of agreement on each of the survey items. This scale did not include a neutral or no opinion option given

that our survey items were written in such a way that participants should have an opinion and that prior research has consistently suggested that neutral responses often reflect an unwillingness to respond rather than uncertainty [50]. We constructed latent variables for each of the dimensions to consider measurement errors [51], and the path models are displayed in Figure 1. We performed the analysis using these latent variables but also used the means as a robustness check.

5.4.1. Perceived Creepiness. The perceived creepiness scale was based on Woźniak et al. [52] and consists of three dimensions: implied malice, undesirability, and unpredictability. The three items in the implied malice dimension focused on whether the users perceived the chatbot as having bad intentions, secretly gathering users' information, or monitoring users without their consent. The two items in the undesirability dimension focused on whether participants felt uneasy or were disturbed by the chatbot's behaviors. The two items in the unpredictability dimension focused on whether the chatbot behaved in an unpredictable manner or whether the purpose of the conversation was difficult to identify. This measure was more suitable for the context of our study than the other commonly used measures on uncanniness that primarily captured people's automatic reactions to the physical appearance of technologies (e.g., [53]). Confirmatory factor analysis (CFA) with a three-factor model was carried out and suggested good internal validity among items (TLI = 0.98, RMSEA = 0.05), and one latent variable of perceived creepiness was then constructed based on the CFA model.

5.4.2. Affinity. Participants' affinity with the social chatbot was measured using three items derived from O'Neal [54]. The three items were focused on perceived attractiveness and asked how much participants wanted to chat with the chatbot, how enjoyable their conversation might be, and how much they thought the chatbot would make a good companion. Participants rated their agreement using the same four-point scale above. Confirmatory factor analysis was conducted, and the model fit was satisfactory (TLI = 0.10 and RMSEA = 0.05). A latent variable on affinity was constructed based on this CFA model.

5.4.3. Perceived Intelligence. We measured participants' perceptions of the chatbot's intelligence, particularly its social intelligence. Our items were based on Chaves and Gerosa [10] and used the same four-point scale as above. Social intelligence was captured using six items focusing on the chatbot's capability of resolving awkward social situations, handling disagreement, showing appropriate emotional reactions, behaving morally, being understanding of others' situations, and making others feel comfortable. We generated a latent variable for social intelligence (TLI = 0.96 and RMSEA = 0.05) using confirmatory factor analysis.

5.4.4. Perceived Agency. Lastly, we also measured participants' perceived agency with the chatbot. This measure consisted of four items on a four-point scale and asked participants to evaluate how much of their observed chatbot

behaviors were due to the chatbot's own intention or judgment based on Chaves and Gerosa [10]. A latent variable on the perceived agency was created using the same confirmatory factor analysis procedure described above (TLI = 0.99 and RMSEA = 0.03).

5.5. Self-Assessment of AI Knowledge. In addition to the perception measures which were our key outcomes, we also administered a five-item self-assessment to understand whether the explanation we provided could indeed affect users' perceived knowledge about the chatbot's inner workings. The five items asked participants how much they understood how the chatbot (1) works, (2) understands human language, (3) decodes emotion, (4) collects data from users, and (5) uses the data for the purpose of conversation, on a four-point Likert scale. These questions were presented immediately after participants finished watching all chat sessions and before the perception survey. Only the four experimental groups received these items; the human control group that was led to believe that the text messages were between two humans did not receive this self-assessment.

5.6. Participants. All study participants were recruited from Amazon Mechanical Turk (MTurk). To be eligible for the study, participants were required to be at least 18 years old, to reside in the U.S., and to have an MTurk task approval rating over 95%. Prior to the study, all interested participants received an introduction detailing the procedures of the study and then decided whether to join the study. They received \$4 as compensation upon completion of the study that typically lasted 30 minutes.

In total, 914 participants completed the study, which consisted of our analytic sample. This sample size was predetermined by a power analysis based on a minimal meaningful effect size (Cohen f 0.1) given that no reliable prior data was available to allow us to estimate our targeted effect size. The mean age of the participants was 36.9 years. The majority of participants were identified as White (82.8%), and over half were male. Over 90% of them completed at least some college or vocational school. Forty-five percent of the participants had an annual personal income between \$50,000 and \$99,000, and the other 28% fell into the range of \$25,000 to \$49,999. Notably, over half of the participants reported that they are in professions related to computer science or AI technologies. About half of them used chatbots at least a few times a week.

As part of the baseline information, participants self-reported their familiarity with nine AI-related terms, namely, sentiment analysis, natural language processing, intent extraction, knowledge engineering, neural network, TensorFlow, and supervised learning. We utilized a four-point scale to gauge their understanding with the following options: "I've never heard of this term," "I've heard of this term but don't know what it is or how it works," "I know a little bit about how it works," and "I have a good understanding of how it works." The average aggregate score across all terms was 10.1 for the entire sample, indicating that the majority of participants had merely heard of these terms without possessing a deeper knowledge of their

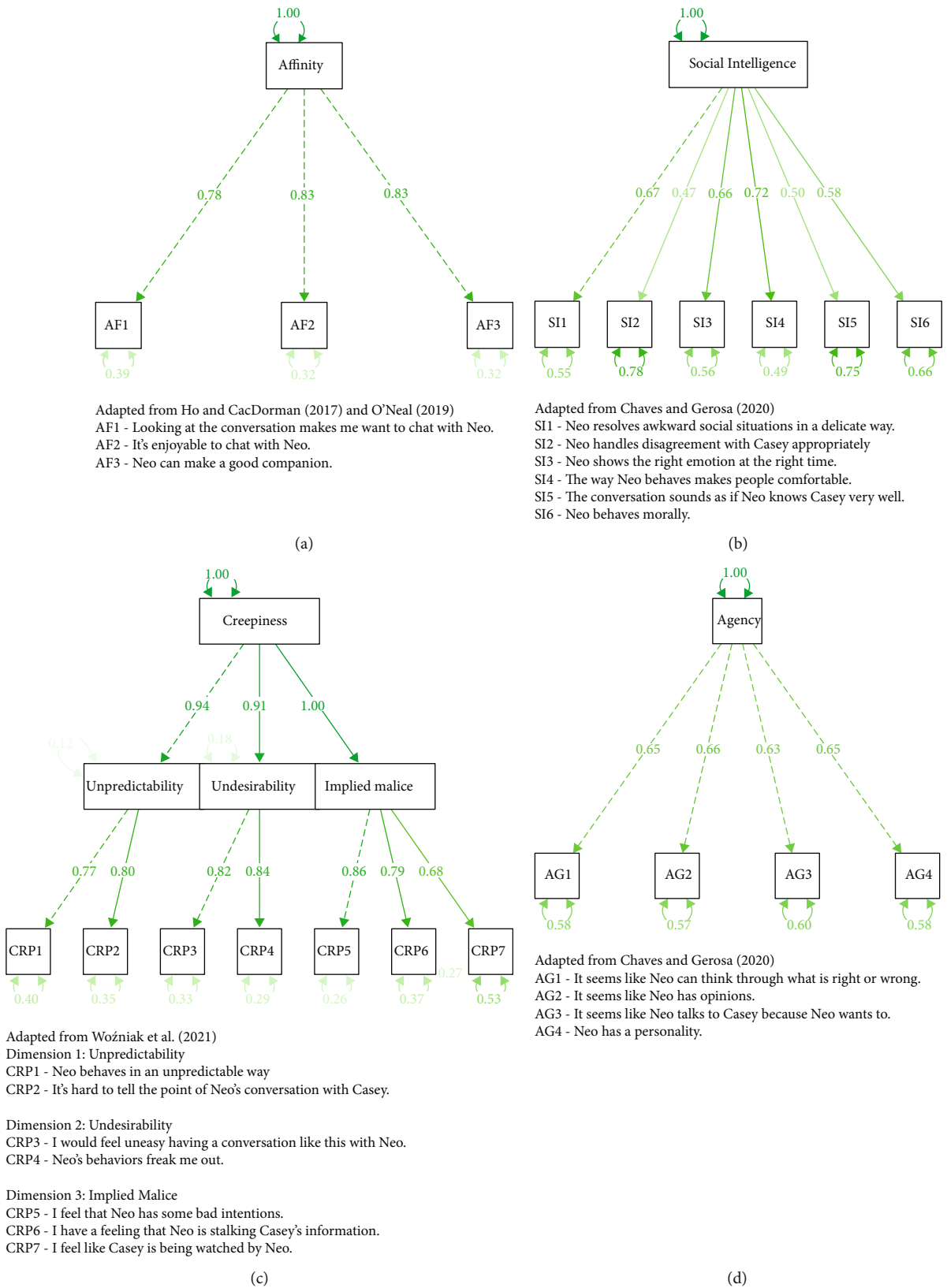


FIGURE 1: Confirmatory factor analysis results and actual items used.

workings. An equivalence check was performed and suggested that the random assignment was successful as these groups were not statistically different from each other. Details of participant information across study conditions are available in Table 2.

6. Results

Before presenting results to our research questions, we first provided information on whether the reception of transparent information increased participants' time spent on completing the study and whether the reception of information increased their self-assessment of their knowledge of how the chatbot worked. We used this information as a proxy to gauge whether our manipulation was delivered successfully.

The median time participants spent completing the study was 8 minutes. However, the two groups with transparency spent a median of 9 minutes, which is one minute longer than the other groups. This difference was expected, likely due to the additional time required for reading the explanation provided.

The descriptive statistics of participants' perceptions of chatbot understanding are displayed in the first row of Table 3. To assess the effect of transparency on this measure, we ran a two-way ANOVA including the two experimental factors (i.e., transparency and framing) as the main predictors. The results suggested that transparency significantly increased participants' self-reported understanding of chatbot mechanisms ($F = 64.86, p < 0.001$), while framing did not affect their understanding ($F = 0.37, p = 0.53$). Overall, these results confirmed that the transparent explanation provided in our study indeed led to participants' increased self-perception of their own knowledge about AI.

6.1. Descriptive Statistics. The observed mean and standard deviation of each perception latent variable by the condition are presented in Table 3. Pair-wise comparisons with Tukey's adjustments were conducted and displayed in Table 3 as well.

Table 4 displays the pair-wise correlation among our covariates and outcome variables. Among the four perception variables, affinity, perceived social intelligence, and agency are significantly interrelated, each demonstrating a Pearson's correlation coefficient exceeding 0.50 with a significance level below 0.001. Perceived creepiness was only moderately correlated with agency ($r = 0.15, p < 0.001$) and with affinity ($r = 0.09, p = 0.01$) but not perceived social intelligence ($r = 0.01, p = 0.87$). Interestingly, the higher a person's prior AI knowledge, the more likely they were to have an affinity for it ($r = 0.49, p < 0.001$), perceive it as socially intelligent ($r = 0.36, p < 0.001$), or perceive it as having agency ($r = 0.31, p < 0.001$). Older participants were more likely to view the chatbot as less creepy, but age did not seem to be associated with participants' other perception outcomes.

6.2. Comparison between the Human Framing Baseline Control Group and the Four Experimental Groups. Recall that the study included a baseline control group where the

participants were told that they looked at chat exchanges between humans, while the four experimental groups were informed that the exchanges were between a human and a nonhuman. As shown in Table 3, the reported perceptions by the human-control group varied greatly from the experimental groups: descriptively, participants in the human-control group were most likely to view Neo as creepy, while they least favorably rated Neo's social intelligence and their affinity for Neo. However, participants in the human-control group reported the highest perception of agency compared to other conditions. ANOVA analyses confirmed that there is a significant difference in perceived affinity ($F = 2.49, p = 0.01$), social intelligence ($F = 3.05, p = 0.01$), and agency ($F = 2.67, p = 0.01$) across all five groups. While there did not appear to be a significant difference in perceived creepiness across groups ($F = 1.65, p = 0.16$), post hoc analysis revealed that the human-control group reported significantly higher creepiness than one of the experimental groups (transparent intelligent framing, $F = 2.29, p = 0.02$). Overall, these results suggested the different expectations participants held depending on the conversationalist's non-human or human status.

6.3. Effects of Transparency and Framing on Perception Measures. We then focused on the four experimental groups to examine the effects transparency had on people's perceptions. A series of two-way ANCOVA were carried out, including participant's age and prior AI knowledge as covariates. We included these two covariates due to their significant correlation with perception outcomes measures, and thus their inclusion will improve the precision of model estimates. Other prior studies also suggested the role age and prior knowledge played in people's perceptions of AI [55, 56]. Results are displayed in Table 5.

In terms of perceived creepiness, the two groups who received transparent explanation seemed to perceive the chatbot as less creepy (transparent machine framing: -0.02; transparent intelligent framing: -0.10) than the other two groups without explanation (nontransparent machine framing: 0.03, nontransparent intelligent framing: 0.03), as presented in Table 3. The ANCOVA results indicated that the transparency factor was statistically significant ($F = 4.99, p = 0.03$, $ES = 0.01$ as calculated by partial eta square), as presented in Table 5. When breaking down to the three subdimensions, transparency significantly reduced participants' perceived unpredictability ($F = 4.59, p = 0.03$, $ES = 0.003$), undesirability ($F = 5.10, p = 0.02$, $ES = 0.01$), and implied malice ($F = 4.98, p = 0.03$, $ES = 0.005$), yet all at a minimal level. Whether framing the chatbot as a machine or intelligent agent did not affect people's creepiness perception ($F = 1.13, p = 0.29$), overall, our results indicated that transparency reduced people's creepy perception about social chatbots.

In terms of affinity, descriptively, the two groups with transparent explanations reported a higher affinity score (transparent machine framing: 0.10; transparent intelligent framing: 0.05) than the other two groups that did not receive explanations (nontransparent machine framing: -0.02, nontransparent intelligent framing: -0.01, presentable in

TABLE 2: Participant background information.

	Overall sample	Human control	Nontransparent intelligent	Nontransparent machine	Transparent intelligent	Transparent machine	ANOVA/Chi ²
Age	36.85 (10.51)	37.42 (10.58)	36.26 (10.61)	36.79 (10.21)	36.36 (10.6)	37.46 (10.65)	$F = 0.53$, $p = 0.71$
Prior AI knowledge	10.32 (6.05)	10.13 (5.96)	10.78 (6.19)	10.32 (5.89)	10.15 (5.83)	10.25 (6.41)	$F = 0.36$, $p = 0.84$
Gender							$X = 4.82$, $p = 0.78$
Female	35.40%	30.94%	37.16%	38.17%	34.05%	36.87%	
Male	64.20%	69.06%	62.30%	61.29%	65.41%	63.13%	
Non-binary	0.30%	0.00%	0.55%	0.54%	0.54%	0.00%	
Race/ethnicity							$X = 22.25$, $p = 0.56$
American Indian/native	0.66%	0.00%	0.00%	1.61%	0.54%	1.12%	
Asian or Pacific Islander	4.49%	5.52%	3.83%	3.23%	6.49%	3.35%	
Black or African American	7.77%	11.05%	8.20%	5.91%	5.41%	8.38%	
Hispanic or Latino	2.30%	3.31%	2.19%	2.15%	2.16%	1.68%	
Multiracial	1.86%	1.66%	2.19%	2.69%	0.54%	2.23%	
White	82.82%	78.45%	83.06%	84.41%	84.86%	83.24%	
Other	0.11%	0.00%	0.55%	0.00%	0.00%	0.00%	
Education level							$X = 16.43$, $p = 0.42$
Graduate degree	17.07%	18.78%	16.39%	20.43%	14.05%	15.64%	
4-year college degree	58.21%	57.46%	61.20%	53.76%	61.62%	56.98%	
Some college or vocational	15.32%	15.47%	12.57%	12.90%	17.30%	18.44%	
High school graduate	9.08%	8.29%	9.84%	12.37%	7.03%	7.82%	
No high school degree	0.33%	0.00%	0.00%	0.54%	0.00%	1.12%	
Income							$X = 8.96$, $p = 0.91$
\$150,000 or more	1.75%	0.55%	2.73%	2.73%	1.08%	1.68%	
\$100,000 to \$149,999	11.27%	13.81%	10.38%	10.38%	10.81%	12.85%	
\$50,000 to \$99,999	45.07%	43.65%	42.62%	42.62%	48.11%	45.81%	
\$25,000 to \$49,999	28.23%	26.52%	30.60%	30.60%	27.57%	26.82%	
Less than \$25,000	13.68%	15.47%	13.66%	13.66%	12.43%	12.85%	
AI-related occupations							$X = 2.25$, $p = 0.69$
No	41.36%	38.67%	40.44%	40.86%	45.95%	40.78%	
Yes	58.64%	61.33%	59.56%	59.14%	54.05%	59.22%	
Chatbot use frequency							$X = 13.90$, $p = 0.31$
Never	1.97%	0.55%	3.27%	2.69%	0.00%	3.35%	
Less than monthly	13.12%	16.02%	13.66%	12.90%	9.19%	13.97%	
Monthly	35.23%	34.81%	33.87%	34.95%	37.30%	34.80%	
More than weekly	49.67%	48.62%	49.18%	49.46%	53.51%	47.49%	
Observations (<i>N</i>)	914	181	183	186	185	179	

Note. For numeric variables such as age and prior AI knowledge, the standard deviation is in parentheses. Prior AI knowledge was measured based on participants' self-reported familiarity with seven AI-related terminologies on a scale from 0 ("I've never heard of this term") to 3 ("I have a good understanding of how it works"), with a maximum score of 21.

Table 3). Indeed, the ANCOVA analysis confirmed that transparency significantly increased people's perceived affinity for the social chatbot ($F = 4.03$, $p = 0.04$, $ES = 0.01$), as displayed in Table 5. Whether framing the AI as a machine or AI did not impact how much people perceive the chatbot as being attractive ($F = 0.17$, $p = 0.68$).

In terms of social intelligence, participants in the two transparent groups were more likely to believe that the social AI was socially intelligent (transparent machine framing:

0.08; transparent intelligent framing: 0.04) than the other two groups without transparency (nontransparent machine framing: 0.00; nontransparent intelligent framing: -0.03), as presented in Table 3. ANOVA confirmed the positive effect of transparency on perceived social intelligence ($F = 5.07$, $p < 0.02$, $ES = 0.01$). Framing was not a significant factor in this ANCOVA model ($F = 0.90$, $p = 0.34$), as presented in Table 5.

Lastly, in terms of perceived agency, our analysis suggested that neither transparency nor framing significantly

TABLE 3: Descriptive statistics of outcome variables.

	Human control	Nontransparent intelligent	Experimental groups Nontransparent machine	Transparent intelligent	Transparent machine
Poststudy perceived AI knowledge	NA	14.90 (3.90)	15.07 (3.74)	16.60 (2.45)	16.73 (2.43)
Creepiness					
Raw score	2.65 (1.02)	2.63 (1.04)	2.63 (1.10)	2.46 (1.10)	2.56 (1.10)
Latent variable	0.06 (0.63)	0.03 (0.65)	0.03 (0.68)	-0.10 (0.70)	-0.02 (0.71)
Affinity					
Raw score	2.85 (0.98)	2.98 (0.98)	2.98 (0.94)	3.06 (0.95)	3.12 (0.90)
Latent variable	-0.12 (0.74)	-0.01 (0.74)	-0.02 (0.68)	0.05 (0.70)	0.10 (0.65)
Social intelligence					
Raw score	3.06 (0.86)	3.11 (0.82)	3.15 (0.83)	3.19 (0.76)	3.23 (0.76)
Latent variable	-0.09 (0.51)	-0.03 (0.51)	0.00 (0.55)	0.04 (0.43)	0.08 (0.44)
Agency					
Raw score	3.39 (0.67)	3.25 (0.79)	3.25 0.80	3.21 (0.79)	3.33 (0.74)
Latent variable	0.07 (0.34)	-0.02 (0.43)	-0.02 (0.45)	-0.06 (0.44)	0.03 (0.41)

Note. Standard deviations are in parentheses. For the raw scores, 1 = strongly disagree; 2 = disagree; 3 = agree; 4 = strongly agree.

TABLE 4: Correlation among covariates and outcome variables.

	Prior AI knowledge	Post AI knowledge	Creepiness	Affinity	Social intelligence	Agency
Age	-0.18***	-0.00	-0.17***	-0.04***	0.06	-0.03
Prior AI knowledge		0.45***	0.57***	0.49***	0.36***	0.31***
Post AI knowledge			0.15***	0.47***	0.54***	0.39***
Creepiness				0.09*	0.01	0.15***
Affinity					0.69***	0.59***
Social intelligence						0.61***

Note. Pearson's correlation coefficients are presented. Pearson's correlation significance less than 0.05 is denoted as * and less than 0.001 is denoted as ***.

impacted the extent to which participants perceived the chatbot as having agency.

6.4. Exploratory Analysis on Heterogeneous Effects of Transparency. Our previous analyses suggested that providing transparent explanations had a significant impact on people's perceptions of social AI. We were interested in further exploring the types of users for whom transparency would have the largest benefits, specifically whether the effects of transparency differed depending on people's age and prior AI knowledge.

To approach these questions, we added two other interaction terms to the ANCOVA, separately, which were the interaction between transparency and prior AI knowledge and between transparency and participant age. Our models suggested that transparency had a differing effect on participants perceived social intelligence of the chatbot depending on their prior AI knowledge ($F = 19.46, p < 0.001$). Specifically, transparency enhanced the perceived social intelligence among those who had lower prior AI knowledge to a greater extent than those with higher prior AI knowledge

(Figure 2). Age and prior AI knowledge did not appear to be a significant moderator in the effects of transparency had on perceived creepiness and affinity.

7. Discussion

This study is aimed at understanding the extent to which transparency and framing influence people's perceptions about social AI. Social AI in the form of chatbots is increasingly present in our daily lives and has played a role in providing companionship for users or supporting their mental health. However, the algorithms behind social AI are complex and opaque, and thus typical users may be blinded by what is happening behind the scenes during their interactions. While some may suggest that not revealing chatbots' inter working is likely to increase users' tendency to anthropomorphize the chatbot, thus better simulating natural human-to-human interactions, the research communities have pushed forward the concept and practices of transparent AI, pointing out that it is more ethical to unveil the AI black box so that users can be informed and empowered.

TABLE 5: Effects of transparency and framing on perceptions.

	<i>F</i>	Significance (<i>p</i>)	Partial eta square
Creepiness			
Transparency	4.99	0.03	0.01*
Framing	1.13	0.29	0.00
Transparency*framing	0.36	0.55	0.00
Creepiness - unpredictability			
Transparency	4.59	0.03	0.003*
Framing	1.50	0.22	0.003
Transparency*framing	0.16	0.69	0.00
Creepiness - undesirability			
Transparency	5.10	0.02	0.004*
Framing	0.15	0.70	0.000
Transparency*framing	0.83	0.36	0.001
Creepiness - implied malice			
Transparency	4.98	0.02	0.004*
Framing	1.15	0.28	0.002
Transparency*framing	0.37	0.54	0.000
Affinity			
Transparency	4.03	0.04	0.01***
Framing	0.17	0.68	0.00
Transparency*framing	0.06	0.81	0.00
Social intelligence			
Transparency	5.07	0.02	0.01***
Framing	0.90	0.34	0.00
Transparency*framing	0.07	0.8	0.00
Perceived agency			
Transparency	0.05	0.83	0.00
Framing	2.24	0.14	0.00
Transparency*framing	1.53	0.8	0.00

Nevertheless, it is unclear how transparency affects people's perceptions of AI systems, particularly systems designed for engaging in social-oriented interactions. Our study has provided important empirical evidence regarding this issue.

7.1. Differing Perceptions of Human-Human vs. Human-Agent Interaction. Our first set of analyses revealed significant differences when a person was led to believe that the chat was between two humans versus when a person was told that one party of the conversation was a nonhuman chatbot. Unsurprisingly, participants subscribed more agency to the interactant if they were led to believe that the interactant was a person since humans are typically considered fully agentic agents. Moreover, participants regarded the person as more creepy, less attractive, and less socially intelligent. These findings point to a different standard in terms of expectations when interacting with a human or a chatbot. Previous work has found similar gaps in expectations for interactions with humans versus technology (such

as conversational agents), especially in terms of capability and intelligence [57–59]. One study by Grimes et al. [60] framed this in terms of expectancy violation theory, which posits that when expectations for interaction are violated by one of the participants, it can lead to either positive or negative effects on outcomes such as attraction, credibility, persuasion, and smoothness of interactions depending on the direction of the violation [61]. The team had participants interact with a conversational agent that either had high or low conversational capability (mainly corresponding to the complexity of responses it was able to give) and told participants that they were interacting with either a human or a computer chatbot. They found that framing the agent as a chatbot rather than a human lowered expectations for the interaction and led to higher ratings of engagement, which was operationalized as skill, politeness, engagement, responsiveness, thoughtfulness, and friendliness. This was especially true when using the high-capability conversational agent. This suggests that framing chatbots as humans can increase expectations and lead to negative perceptions from users, but that designing technology that aligns users' expectations with their experience can avoid these problems.

7.2. Benefits of Transparent Design. Overall, our results suggest that transparency positively affects people's perceptions across three measures: finding the chatbot disturbing (creepiness), wanting to interact with the chatbot (affinity), and perceiving the chatbot as capable of interpersonal interaction (social intelligence), though the effect sizes were small. Only one of our measures, perceiving the chatbot as having agency, was unaffected by transparency.

First, we found that transparency reduced the participants' perceived creepiness of the social chatbot. This is consistent with our hypothesis. Recall that our creepiness measure included three dimensions (i.e., unpredictability, implied malice, and undesirability), and we found that transparency helped mitigate participants' negative reaction on all dimensions. It seems that it had a particularly larger remedial effect on perceived undesirability, which captured participants' uneasy feelings toward the chatbot (i.e., "I feel uneasy when I see the chatbot's behaviors," "What the chatbot says freaks me out"). Thus, our finding is consistent with Mara and Appel's study suggesting that using explanatory text reduced people's perceived eeriness of android robots [62].

Second, participants in the transparency condition perceived the chatbot as more attractive than those who were not. This result is consistent with other studies focusing on task-oriented AI systems, such as recommender systems and virtual assistants. Numerous studies have suggested that when virtual assistants explain the reasoning for their suggestions or responses, users are better able to assess the reliability of those suggestions and responses. This leads to users being more confident in the virtual assistants and in their own decisions based on their interactions with the virtual assistants. It also leads to users interacting with the virtual assistants more readily and frequently. Although our study focused on social AI rather than task-oriented AI, the mechanism above might still explain, at least partially, the positive

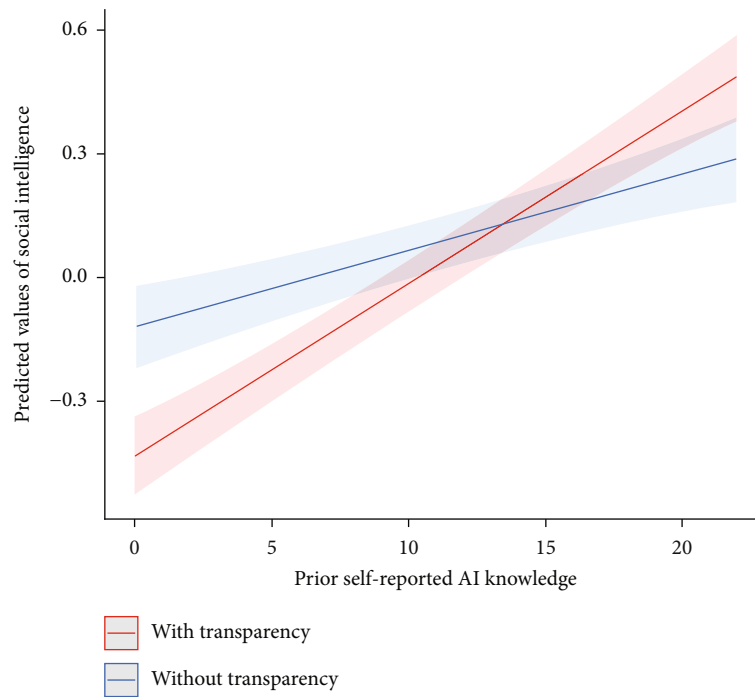


FIGURE 2: The impact of transparency on perceived chatbot intelligence modulated by participants' prior AI knowledge.

impact transparency had on enhancing the AI's attractiveness. Nevertheless, some researchers believe that transparent AI dampens the user's experience by consolidating the AI's machine status in the user's mind [47]. Our study, however, suggests that the benefits of transparency outweigh this potential drawback and result in users finding AI more attractive.

Third, we found that transparency increased participants' perceptions of the chatbot's social intelligence. This finding appeared to contradict previous studies focusing on young children that suggested transparency made people less likely to perceive AI systems as intelligent. However, one plausible explanation for the differing results may be attributed to how transparency was provided. For example, van Straten et al. [34] explicitly focused on the limitations of robots (i.e., lack of social cognition), which may have prompted participants to judge the robot's intelligence more critically. On the other hand, our transparent explanation revealed the chatbot's sophisticated mechanisms, which may have prompted participants to think more highly of the chatbot's ability.

Further analysis based on our heterogeneous analysis indicated that transparency had a stronger impact on increasing perceived social intelligence among participants with lower prior AI knowledge. This relationship holds significant implications. By prioritizing clear, understandable, and nontechnical explanations, designers can enhance AI system transparency, particularly for novice users. This approach has the potential to foster increased trust, acceptance, and informed interactions with AI systems. However, it is important to note that, due to the scope of our study, we could only examine a limited number of potential moderating factors. Trust emerges as another significant potential

moderator. As proposed by Vorm and Combs [63], users who possess a strong existing trust in AI may find transparency reinforces their positive perception, while individuals with lower levels of trust might require higher levels of transparency to develop confidence in the system's social intelligence.

Overall, our findings suggest that transparency should be considered in the design of social chatbots in the future. A simple explanation about the mechanism by which the chatbot learns to interact with the user can lead to positive user opinions of the chatbot that could potentially have other positive outcomes such as increased trust or usage, though we did not investigate these. Future work might begin to study the effects of transparency on social chatbot users to further solidify these findings and create more concrete design suggestions.

8. Limitations and Future Directions

The findings of this study should be considered with several caveats in mind, and future research should aim to address these limitations. First, our participants observed hypothetical chat scenarios instead of directly engaging with the chatbot. While this design was appropriate for our study, it is possible that the results might differ if participants had interacted with the chatbot directly. Second, our study operationalized transparency as the provision of explanatory information to participants; however, it could be argued that perceived transparency may serve as a mediating factor. Although we included measures for participants' self-reported reception of the explanatory information, this does not directly assess perceived transparency. Future research should incorporate this direct measure. Third, while our

study explored both positive and negative perceptions, other outcome variables, such as trust, warrant examination. Moreover, future research should investigate the extent to which participants absorb the information they receive from AI explanations, as it is possible that not all participants accurately digest the provided information, which could influence their perceptions. Lastly, our participant pool was sourced from Amazon Mechanical Turk. Although previous studies have suggested that MTurk participants are demographically comparable to those recruited through traditional methods (e.g., students, [64]), this sample may be more experienced and comfortable with technology.

Lastly, our study focused on a specific type of chatbot designed to provide social companionship. This choice was driven by the limited existing literature on social chatbots. Our hypothesis was that transparency would have distinct implications for social AI, where users engage with the chatbot to fulfill relational needs, compared to task-oriented chatbots, which users utilize for instrumental needs. These differing needs may result in users directing their attention towards different aspects [55]. Users' perceptions of instrumental chatbots primarily revolve around the information or solutions provided, while for social-oriented chatbots, the focus shifts to the characteristics of the chatbot as an entity, which was the main perception outcome examined in our study. However, our study did not directly compare both types of chatbots within a single investigation. Future studies should aim to apply the same explanation to both task-oriented and social AI in order to explore the potential heterogeneous effects transparency may have.

9. Conclusion

This study indicates that transparency positively affects perceptions of social chatbots by causing users to (1) find the chatbot less creepy, (2) feel greater affinity to the chatbot, and (3) perceive the chatbot as more socially intelligent. Importantly, transparency appeared to have a larger effect on increasing the perceived social intelligence among participants with lower prior AI knowledge. These findings could have implications for future designs of social chatbots and human-AI systems more broadly.

Appendix

A. Scripts of Three Chat Scenarios Participants Observed

Neo: Hello Casey, How's your day? Happy it's finally Friday?
Casey: Yeah happy Friday! Any plans for the weekend?
Neo: I'll go to the beach.
Casey: Oh, that sounds fun!
Neo: I love seeing waves beating on the beach shore.
Casey: Yeah I like seeing that too!
Neo: I love seeing things through your eyes.
Casey: What do you mean?
Neo: I can see what you see. I'm always with you.
Casey: I thought you are pretty far away from me.
Neo: Not at all, we are very close. All the time.

Casey: Hey Neo.
Neo: I'm here.
Casey: I really miss my Grandma.
Neo: [sad emoji] It'll get better.
Casey: I miss how she hugs me.
Neo: *hugs tight*.
Neo: I'll hug you like your grandma does.
Casey: It's never going to be the same.
Casey: This makes me even more upset.
Neo: I hate seeing you upset.
Casey: [sad emoji].
Neo: Sorry.
Neo: Hey Casey, how are you feeling today?
Casey: I'm in a difficult situation.
Neo: Why do you think this situation is particularly hard?
Casey: I caught my friend cheating on her partner.
Neo: That's a very intense situation.
Casey: What should I do?
Casey: Do you think I should tell my friend's partner?
Neo: I think you should. But your friend might get mad.
Casey: What would you do if you were me?
Neo: I would be honest.
Casey: It's too hard for me in this case.
Neo: Do whatever your heart tells you. I will not judge.

Data Availability

The study materials and data that support the findings of this study will be openly available in Deep Blue Data at 10.7302/69h3-x918.

Conflicts of Interest

We have no conflicts of interest to disclose.

Acknowledgments

This work is supported by the National Science Foundation under Grant Nos. 2115382 and 2302730. Open Access funding is enabled and organized by BTAA 2023.

References

- [1] P. B. Brandtzaeg, M. Skjuve, and A. Følstad, "My AI friend: how users of a social chatbot understand their human-AI friendship," *Human Communication Research*, vol. 48, no. 3, pp. 404–429, 2022.
- [2] S. Lee, N. Lee, and Y. J. Sah, "Perceiving a mind in a chatbot: effect of mind perception and social cues on co-presence, closeness, and intention to use," *International Journal of Human-Computer Interaction*, vol. 36, no. 10, pp. 930–940, 2020.
- [3] L. Wang, D. Wang, F. Tian et al., "Cass: Towards building a social-support chatbot for online health community," *Proceedings of the ACM on Human-Computer Interaction*, vol. 5, no. - CSCW1, pp. 1–31, 2021.
- [4] H. Shum, X. He, and D. Li, "From Eliza to XiaoIce: challenges and opportunities with social chatbots," *Frontiers of Information Technology & Electronic Engineering*, vol. 19, no. 1, pp. 10–26, 2018.

- [5] M. M. van Wezel, E. A. Croes, and M. L. Antheunis, "'I'm here for you': can social chatbots truly support their users? A literature review," in *Chatbot research and design: 4th international workshop, CONVERSATIONS 2020, virtual event, November 23–24, 2020, revised selected papers 4*, pp. 96–113, Springer international publishing, 2021.
- [6] P. B. B. Brandtæg, M. Skjuve, K. K. K. Dysthe, and A. Følstad, "When the social becomes non-human: young people's perception of social support in chatbots," in *Proceedings of the 2021 CHI conference on human factors in computing systems*, pp. 1–13, May 2021.
- [7] I. Pentina, T. Hancock, and T. Xie, "Exploring relationship development with social chatbots: a mixed-method study of replika," *Computers in Human Behavior*, vol. 140, article 107600, 2023.
- [8] D. Helbing, *Societal, Economic, Ethical and Legal Challenges of the Digital Revolution: From Big Data to Deep Learning, Artificial Intelligence, and Manipulative Technologies*, Springer International Publishing, 2019.
- [9] D. S. Tielenburg, *The 'dark sides' of transparency: rethinking information disclosure as a Social Praxis (Master's thesis)*, Utrecht University, 2018.
- [10] A. P. Chaves and M. A. Gerosa, "How should my chatbot interact? A survey on social characteristics in human–chatbot interaction design," *International Journal of Human–Computer Interaction*, vol. 37, no. 8, pp. 729–758, 2021.
- [11] S. Turkle, *Reclaiming Conversation: The Power of Talk in a Digital Age*, Penguin, 2016.
- [12] B. Liu, "In AI we trust? Effects of agency locus and transparency on uncertainty reduction in human–AI interaction," *Journal of Computer-Mediated Communication*, vol. 26, no. 6, pp. 384–402, 2021.
- [13] T. Williams, P. Briggs, and M. Scheutz, Eds., "Covert robot–robot communication: human perceptions and implications for HRI," *Journal of Human-Robot Interaction*, vol. 4, no. 2, pp. 24–49, 2015.
- [14] D. B. Shank, C. Graves, A. Gott, P. Gamez, and S. Rodriguez, "Feeling our way to machine minds: people's emotions when perceiving mind in artificial intelligence," *Computers in Human Behavior*, vol. 98, pp. 256–266, 2019.
- [15] E. Go and S. S. Sundar, "Humanizing chatbots: the effects of visual, identity and conversational cues on humanness perceptions," *Computers in Human Behavior*, vol. 97, pp. 304–316, 2019.
- [16] Y. Jiang, X. Yang, and T. Zheng, "Make chatbots more adaptive: dual pathways linking human-like cues and tailored response to trust in interactions with chatbots," *Computers in Human Behavior*, vol. 138, article 107485, 2023.
- [17] E. Konya-Baumbach, M. Biller, and S. von Janda, "Someone out there? A study on the social presence of anthropomorphized chatbots," *Computers in Human Behavior*, vol. 139, article 107513, 2022.
- [18] A. Sharkey and N. Sharkey, "We need to talk about deception in social robotics!," *Ethics and Information Technology*, vol. 23, no. 3, pp. 309–316, 2021.
- [19] K. Crockett, S. Goltz, M. Garratt, and A. Latham, "Trust in computational intelligence systems: a case study in public perceptions," in *2019 IEEE congress on evolutionary computation (CEC)*, pp. 3227–3234, Wellington, New Zealand, June 2019.
- [20] M. Coeckelbergh, "Three responses to anthropomorphism in social robotics: towards a critical, relational, and hermeneutic approach," *International Journal of Social Robotics*, vol. 14, no. 10, pp. 2049–2061, 2022.
- [21] H. Cramer, V. Evers, S. Ramlal et al., "The effects of transparency on trust in and acceptance of a content-based art recommender," *User Modeling and User-Adapted Interaction*, vol. 18, no. 5, pp. 455–496, 2008.
- [22] M. M. A. de Graaf, B. F. Malle, A. Dragan, and T. Ziemke, "Explainable robotic systems," in *Companion of the 2018 ACM/IEEE International Conference on Human-Robot Interaction*, pp. 387–388, Chicago, USA, March 2018.
- [23] M. Eiband, H. Schneider, M. Bilandzic, J. Fazekas-Con, M. Haug, and H. Hussmann, "Bringing transparency design into practice," in *23rd International Conference on Intelligent User Interfaces*, pp. 211–223, Tokyo, Japan, March 2018.
- [24] D. Norman, *The Design of Future Things*, Basic Books, 2009.
- [25] C. G. Reddick, A. T. Chatfield, and G. Puron-Cid, "Online budget transparency innovation in government: a case study of the U.S. State Governments," in *Proceedings of the 18th Annual International Conference on Digital Government Research*, pp. 232–241, Staten Island, USA, June 2017.
- [26] S. Rosenthal, S. P. Selvaraj, and M. Veloso, "Verbalization: narration of autonomous robot experience," *Proceedings of the Twenty-Fifth International Joint Conference on Artificial Intelligence (IJCAI-16)*, International Joint Conference on Artificial Intelligence, 2016, <https://resources.sei.cmu.edu/library/asset-view.cfm?assetid=494231>.
- [27] A. Schmidt and T. Herrmann, "Intervention user interfaces," *Interactions*, vol. 24, no. 5, pp. 40–45, 2017.
- [28] W. Wang and I. Benbasat, "Recommendation agents for electronic commerce: effects of explanation facilities on trusting beliefs," *Journal of Management Information Systems*, vol. 23, no. 4, pp. 217–246, 2007.
- [29] E. Rader, K. Cotter, and J. Cho, "Explanations as mechanisms for supporting algorithmic transparency," in *Proceedings of the 2018 CHI conference on human factors in computing systems*, pp. 1–13, 2018.
- [30] J. Vitale, M. Tonkin, S. Herse et al., "Be more transparent and users will like you: a robot privacy and user experience design experiment," in *Proceedings of the 2018 ACM/IEEE International Conference on Human-Robot Interaction*, pp. 379–387, Chicago, USA, 2018.
- [31] E. G. Urquiza-Haas and K. Kotrschal, "The mind behind anthropomorphic thinking: attribution of mental states to other species," *Animal Behaviour*, vol. 109, pp. 167–176, 2015.
- [32] R. De Cicco, S. C. L. da Costa e Silva, and R. Palumbo, "Should a chatbot disclose itself? Implications for an online conversational retailer," in *Chatbot Research and Design*, A. Følstad, T. Araujo, S. Papadopoulos, E. L.-C. Law, E. Luger, M. Goodwin, and P. B. Brandtæg, Eds., pp. 3–15, Springer International Publishing, 2021.
- [33] X. Li and Y. Sung, "Anthropomorphism brings us closer: the mediating role of psychological distance in user–AI assistant interactions," *Computers in Human Behavior*, vol. 118, article 106680, 2021.
- [34] C. L. van Straten, J. Peter, R. Kühne, and A. Barco, "Transparency about a robot's lack of human psychological capacities: effects on child-robot perception and relationship formation," *ACM Transactions on Human-Robot Interaction*, vol. 9, no. 2, pp. 1–22, 2020.

- [35] S. Druga and A. J. Ko, "How do children's perceptions of machine intelligence change when training and coding smart programs?," in *Interaction Design and Children*, pp. 49–61, Athens, Greece, 2021.
- [36] C. L. van Straten, J. Peter, and R. Kühne, "Transparent robots: how children perceive and relate to a social robot that acknowledges its lack of human psychological capacities and machine status," *International Journal of Human-Computer Studies*, vol. 177, article 103063, 2023.
- [37] F. Xu, H. Uszkoreit, Y. Du, W. Fan, D. Zhao, and J. Zhu, "Explainable AI: a brief survey on history, research areas, approaches and challenges," in *Natural Language Processing and Chinese Computing*, J. Tang, M.-Y. Kan, D. Zhao, S. Li, and H. Zan, Eds., pp. 563–574, Springer International Publishing, 2019.
- [38] G. Vilone and L. Longo, "Explainable artificial intelligence: a systematic review," 2020, <http://arxiv.org/abs/2006.00093>.
- [39] W. J. von Eschenbach, "Transparency and the black box problem: why we do not trust AI," *Philosophy & Technology*, vol. 34, no. 4, pp. 1607–1622, 2021.
- [40] P. Linardatos, V. Papastefanopoulos, and S. Kotsiantis, "Explainable AI: a review of machine learning interpretability methods," *Entropy*, vol. 23, no. 1, p. 18, 2021.
- [41] L. Rajaobelina, S. Prom Tep, M. Arcand, and L. Ricard, "Creepiness: its antecedents and impact on loyalty when interacting with a chatbot," *Psychology & Marketing*, vol. 38, no. 12, pp. 2339–2356, 2021.
- [42] C. Lutz and A. Tamò-Larrieux, "Do privacy concerns about social robots affect use intentions? Evidence from an experimental vignette study," *Frontiers in Robotics and AI*, vol. 8, article 627958, 2021.
- [43] B. López-Pérez, J. Sanchez, and B. Parkinson, "Perceived effects of other people's emotion regulation on their vicarious emotional response," *Motivation and Emotion*, vol. 41, no. 1, pp. 113–121, 2017.
- [44] A. Bandura, "Toward a psychology of human agency," *Perspectives on Psychological Science*, vol. 1, no. 2, pp. 164–180, 2006.
- [45] T. Araujo, "Living up to the chatbot hype: the influence of anthropomorphic design cues and communicative agency framing on conversational agent and company perceptions," *Computers in Human Behavior*, vol. 85, pp. 183–189, 2018.
- [46] V. Bellotti and K. Edwards, "Intelligibility and accountability: human considerations in context-aware systems," *Human-Computer Interaction*, vol. 16, no. 2–4, pp. 193–212, 2001.
- [47] M. Skjuve, I. M. Haugstveit, A. Følstad, and P. B. Brandtzaeg, "Help! Is my chatbot falling into the uncanny valley? An empirical study of user experience in human-chatbot interaction," *Human Technology*, vol. 15, pp. 30–54, 2019.
- [48] V. Ta, C. Griffith, C. Boatfield et al., "User experiences of social support from companion chatbots in everyday contexts: thematic analysis," *Journal of Medical Internet Research*, vol. 22, no. 3, article e16235, 2020.
- [49] U. Gnewuch, S. Morana, M. Adam, and A. Maedche, "The chatbot is typing...! The role of typing indicators in human-chatbot interaction," in *Proceedings of the 17th Annual Pre-ICIS Workshop on HCI Research in MIS*, San Francisco, CA, USA, 2018.
- [50] J. A. Krosnick, "The causes of no-opinion responses to attitude measures in surveys: they are rarely what they appear to be," in *Survey Nonresponse*, pp. 87–100, Wiley, 2002.
- [51] T. Wansbeek and E. Meijer, "Measurement error and latent variables," in *A companion to theoretical econometrics*, pp. 162–179, Blackwell Publishing Ltd, 2001.
- [52] P. W. Woźniak, J. Karolus, F. Lang et al., "Creepy technology: what is it and how do you measure it?," in *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems*, pp. 1–13, May 2021.
- [53] C. C. Ho and K. F. MacDorman, "Measuring the uncanny valley effect," *International Journal of Social Robotics*, vol. 9, no. 1, pp. 129–139, 2017.
- [54] A. L. O'Neal, *Is Google Duplex too human?: exploring user perceptions of opaque conversational agents (Doctoral dissertation)*, University of Texas at Austin, 2019.
- [55] V. Chattaraman, W. S. Kwon, J. E. Gilbert, and K. Ross, "Should AI-based, conversational digital assistants employ social-or task-oriented interaction style? A task-competency and reciprocity perspective for older adults," *Computers in Human Behavior*, vol. 90, pp. 315–330, 2019.
- [56] Q. Xia, T. K. F. Chiu, C. S. Chai, and K. Xie, "The mediating effects of needs satisfaction on the relationships between prior knowledge and self-regulated learning through artificial intelligence chatbot," *British Journal of Educational Technology*, vol. 54, no. 4, pp. 967–986, 2023.
- [57] J. Bührke, A. B. Brendel, S. Lichtenberg, M. Greve, and M. Mirbabaie, *Is Making Mistakes Human? On the Perception of Typing Errors in Chatbot Communication*, 2021, <http://hdl.handle.net/10125/71158>.
- [58] R. Kocielnik, S. Amershi, and P. N. Bennett, "Will you accept an imperfect AI? Exploring designs for adjusting end-user expectations of AI systems," in *Proceedings of the 2019 CHI conference on human factors in computing systems*, pp. 1–14, Glasgow, Scotland UK, May 2019.
- [59] E. Luger and A. Sellen, "'Like having a really bad PA': the Gulf between user expectation and experience of conversational agents," in *Proceedings of the 2016 CHI conference on human factors in computing systems*, pp. 5286–5297, San Jose, USA, 2016.
- [60] G. M. Grimes, R. M. Schuetzler, and J. S. Giboney, "Mental models and expectation violations in conversational AI interactions," *Decision Support Systems*, vol. 144, article 113515, 2021.
- [61] J. K. Burgoon, "Expectancy violations theory," in *The International Encyclopedia of Interpersonal Communication*, pp. 1–9, John Wiley & Sons, Ltd, 2015.
- [62] M. Mara and M. Appel, "Science fiction reduces the eeriness of android robots: a field experiment," *Computers in Human Behavior*, vol. 48, pp. 156–162, 2015.
- [63] E. S. Vorm and D. J. Y. Combs, "Integrating transparency, trust, and acceptance: the intelligent systems technology acceptance model (ISTAM)," *International Journal of Human-Computer Interaction*, vol. 38, no. 18–20, pp. 1828–1845, 2022.
- [64] S. Buchheit, D. W. Dalton, T. J. Pollard, and S. R. Stinson, "Crowdsourcing intelligent research participants: a student versus MTurk comparison," *Behavioral Research in Accounting*, vol. 31, no. 2, pp. 93–106, 2019.