

Anthropomorphism Moderates the Relationships of Dispositional, Perceptual, and Behavioral Trust in a Robot Teammate

Proceedings of the Human Factors and Ergonomics Society Annual Meeting 2023, Vol. 67(1) 529–536
Copyright © 2023 Human Factors and Ergonomics Society
DOI: 10.1177/21695067231196240
journals.sagepub.com/home/pro



Myke C. Cohen¹, Matthew A. Peel¹, Matthew J. Scalia¹,
Matthew M. Willett¹, Erin K. Chiou¹, Jamie C. Gorman¹,
and Nancy J. Cooke¹

Abstract

Trust plays a critical role in the success of human-robot teams (HRTs). While typically studied as a perceptual attitude, trust also encompasses individual dispositions and interactive behaviors like compliance. Anthropomorphism, the attribution of human-like qualities to robots, is a related phenomenon that designers often leverage to positively influence trust. However, the relationship of anthropomorphism to perceptual, dispositional, and behavioral trust is not fully understood. This study explores how anthropomorphism moderates these relationships in a virtual urban search and rescue HRT scenario. Our findings indicate that the moderating effects of anthropomorphism depend on how a robot's recommendations and its confidence in them are communicated through text and graphical information. These results highlight the complexity of the relationships between anthropomorphism, trust, and the social conveyance of information in designing for safe and effective human-robot teaming.

Keywords

Human-Robot Teaming, Anthropomorphism, Trust, Communication, Dispositional Trust, Perceptual Trust, Behavioral Trust, Confidence

Introduction

Robots have been deployed in the field for applications in which human lives hang in the balance, such as urban search and rescue (USAR; Casper & Murphy, 2003). Today, advances in artificial intelligence (AI) have brought us closer to the reality of human-robot teams (HRTs), in which people and robots work interdependently toward mutual goals, resembling teamwork in human teams. Effective interactions are crucial to the success and safety of increasingly complex HRTs; as such, there is a pressing need to understand the factors that influence the effectiveness of HRT interactions (Cooke et al., 2023).

Even as operators or supervisors, people tend to interact with non-humans like robots in ways that can be described through human social norms and constructs (Epley et al., 2007; Nass & Moon, 2000). The conformity of human-robot interactions to social norms like politeness influences how people trust their robot counterparts, including their reliance on robot inputs and processes when appropriate (Nass, 2004; Parasuraman & Miller, 2004). Trust has been defined as the

attitude that a robot will act beneficially toward one's goals in risk-prone situations (Lee & See, 2004), including the willingness to be vulnerable to it (Mayer et al., 1995). Because trust is formed, expressed, and updated through repeated interactions, it is crucial to the long-term effectiveness of HRTs (Chiou & Lee, 2023)—particularly as people may be even more likely to rely on human social expectations when interacting with complex robots that are capable of teaming (Groom & Nass, 2007). However, the current understanding of trust is convoluted by its varied conceptualizations and measurements in the literature.

Most studies measure people's *perceptual trust* in an agent, often through self-report questionnaires about its trustworthiness after a period of interaction (Kohn et al., 2021). Another

¹Human Systems Engineering, Arizona State University, Mesa, AZ, USA

Corresponding Author:

Myke C. Cohen, Human Systems Engineering, Arizona State University,
7001 E Williams Field Rd, Mesa, AZ 85212, USA.

Email: myke.cohen@asu.edu

trust construct typically measured through questionnaires before interactions take place is *dispositional trust*: a stable trait that describes people's propensity to trust a robot (Jessup et al., 2019). Some argue, though, that a person's trust in a robot is unambiguous only when perceptual and dispositional trust translate to observable *behavioral trust*, such as compliance with recommended decisions (Meyer & Lee, 2013). Although usually correlated, these measures sometimes offer divergent trust narratives (Hancock et al., 2011). This is not necessarily undesirable—a person who generally perceives a robot as trustworthy ideally complies only with appropriate recommendations. A recent review by Kohn et al. (2021) suggests that designing for trustworthy robot teammates must be based on models that integrate the various measures for these unique but overlapping trust constructs. We posit that as a precursor to this multi-measure approach, we must first establish how dispositional, perceptual, and behavioral trust are interrelated with factors that are used to influence trust, such as anthropomorphism.

Anthropomorphism is the imputation of human traits and qualities to non-human entities (Epley et al., 2007). More anthropomorphic perceptions of an agent form as people socially perceive and interact with it, thus also generally coinciding with greater trust (Waytz et al., 2014). This relationship is the basis for some theoretical frameworks for maintaining trust in HRTs using robot explanations, apologies, and blame redirection (e.g., de Visser et al., 2020). However, the role of anthropomorphism in how robot socializations affect trust is poorly understood. This is partially because current research tends to conflate human-like designs with anthropomorphism—which, like trust, is a complex cognitive phenomenon influenced not only by an agent's design characteristics but also by individual dispositions to perceive robots socially (Fischer, 2011).

To illustrate, Kulms and Kopp (2019) reported that more human-like designs of a virtual advisor's appearance improve perceptual trust but do not affect behavioral trust, and that conversely, the quality of its advice impacts only behavioral trust but not perceptual trust. Such findings cannot be readily ascribed to anthropomorphism, because the presence of humanlike visual features does not guarantee that an agent will be anthropomorphized or trusted as intended (Mori, 1970). People are more likely to anthropomorphize a non-human when they think that it seems capable of human-like thought (Epley et al., 2007). Such opinions tend to be informed more by the humanlikeness of social interactions than by visual appearances (Stein & Ohler, 2017). Indeed, Jensen et al. (2020) showed that machine-like and human-like communication styles can result in different levels of perceived anthropomorphism, though not necessarily perceptual or behavioral trust. On the other hand, de Visser et al. (2016) found that increasing the human-likeness of a virtual agent, including how it gives feedback about its reliability, can: (1) result in different levels of anthropomorphism, perceptual trust, and behavioral

trust; and (2) minimize the magnitudes by which errors decreased trust in both forms. Nevertheless, it remains uncertain which human-like robot socializations concurrently impact trust and anthropomorphism.

Supposing that trusting and anthropomorphic attitudes toward robot teammates form interdependently (M. C. Cohen et al., 2021), predicting how a robot's communication abilities will affect trust must also account for how they may contribute to it being perceived anthropomorphically. In this study, we explore two questions surrounding how communication-related anthropomorphism moderates the relationships between dispositional, perceptual, and behavioral trust. First, do more anthropomorphic perceptions of a robot result in more positive correlations between perceptual and behavioral trust? Second, does a person's perceived anthropomorphism affect the translation of their trusting dispositions into perceived or behavioral trust?

Method

We explore how anthropomorphism moderates trust in a simulated USAR HRT. In this study, the style and presence of confidence indicators in robot communication were manipulated in a 3 x 2 mixed design. *Communication Style* was a between-subjects variable, with communication presented either Graphics-only, Text-only, or a "Full" combination of both. *Robot Confidence*—a robot's conveyed confidence in its own advice—was a within-subjects variable with two levels over two missions: Confidence Displayed or Confidence Absent.

Due to the anthropomorphism questionnaire being administered only at the end of the first mission, we consider only the first mission in this analysis, with Robot Confidence treated as a between-subjects condition. We hypothesized that, across all conditions, anthropomorphism moderates how perceptual trust predicts behavioral trust (H1); and how dispositional trust predicts perceptual trust (H2) and behavioral trust (H3).

Participants

A total of 66 participants were recruited from Arizona State University and online student message boards; 56 were between 18-25 years old, eight between 26-35, and two between 36-55. There were 31 women, 34 men, and one who did not disclose their gender. All participants had normal or corrected-to-normal vision and spoke fluent English. They each received a \$30 Amazon gift card as compensation for participating in the 2-hour study. Participants were randomly assigned to one of the three Communication Style conditions and subjected to a counterbalanced mission order for the two Robot Confidence conditions ($n = 11$). Because of an error in data collection, there were $n = 13$ participants in the Text-only, Confidence Displayed condition, and $n = 9$ for the Text-only, Confidence Absent condition.

Roblox USAR Human-Robot Teaming Testbed

The experiment was conducted remotely using Zoom and a synthetic task environment (STE; Cooke & Shope, 2004) developed in Roblox. The STE was designed to simulate a USAR task in a hotel that collapsed due to an environmental disaster, with survivors dispersed across two floors. Detailed information about the testbed is available in Raimondo et al. (2022).

Participants were told that they were remotely paired with an autonomous USAR robot to search the collapsed structure for survivors. In reality, the robot was controlled by a trained experimenter through a Wizard of Oz technique (Riek, 2012). The robot teammate performed basic navigation tasks within the game environment on its own, such as obstacle detection and avoidance, environmental scanning, and survivor detection. Participants were tasked with issuing and monitoring the execution of high-level navigational commands (e.g., directing the robot to search certain map areas). They interacted with the robot through an interface consisting of a first-person camera view of the robot's movements, a live map of the structure highlighting the robot's current location, a mission timer, a resource panel, and a text chat interface used for navigation commands.

Upon discovering a survivor, the robot offered preliminary assessments of their health status and level of injuries, along with suggestions on which medical resources are needed for treatment. Depending on the experimental condition, these recommendations were communicated through graphical logos, textual explanations, or combinations of both (Figure 1). Confidence indicators, when present, were shown graphically through horizontal bars and textually through percentages. Hazards in the environment, such as active fires, gas leaks, limited visibility, and collapsed passageways, were also present and affected the robot's ability to make accurate recommendations. For all participants, the robot's recommendation accuracy was 70%; all inaccurate recommendations were about the same survivors. Participants were then responsible for verbally reporting the survivor's location to the experimenter, including which medical resource should be used out of three available options based on their own assessment of the survivor's severity condition. The outcome of the participants' resource allocation was evaluated by the experimenter, who provided feedback on whether the victims were rescued successfully or not.

Procedure

Participants joined their scheduled Zoom session, were given informed consent on what would be expected of them in the study. An initial questionnaire was administered, including demographics and propensity to trust automation. Participants then viewed a 15-minute training video describing their mission and the task environment, before completing a 5-minute hands-on training mission. After training, participants completed two 20-minute missions, which were



Figure 1. Communication styles: Graphics-only (left) and Text-only (right). Participants in the Full condition saw both styles in the above configuration.

each followed by trust questionnaires. After Mission 1 participants also completed an anthropomorphism survey. At the end of the experiment, participants were debriefed and compensated for their time.

Measures

Dispositional trust was measured using the Propensity to Trust Automation Scale (Jessup et al., 2019). The scale has six questions measured on a seven-point Likert scale. Examples include “Technology is reliable” and “I rely on technology”.

Perceptual trust was measured using an adapted version of the Chancey et al. (2017) trust questionnaire, consisting of 15 items rated on a seven-point Likert scale. The items were adapted to reference the “robot” that participants interacted with in the study. An example is “The robot always provides the advice I require to help me perform well”.

Behavioral trust was defined in this study as a participant's binary compliance with the robot's triage recommendations, following Meyer and Lee (2013). We measured this as the ratio of the number of times the participant adopted the robot's recommendations to the number of survivors the participant found.

Anthropomorphism was measured using the Godspeed questionnaire (Bartneck et al., 2009). The Godspeed questionnaire comprises 25 items that measure social perceptions about an autonomous agent; responses to the first five questions measure anthropomorphism and were averaged for this analysis. Following Bartneck et al. (2009), we administered this as a fivepoint semantic differential scale. An example includes comparing whether the agent behaved “machine-like” versus “humanlike”. Higher scores indicated more anthropomorphic perceptions of the agent ($M = 2.78$, $SD = 0.98$).

Results

Anthropomorphism was tested as a moderator for the relationships between (a) perceptual and behavioral trust; (b) dispositional and perceptual trust; and (c) dispositional and

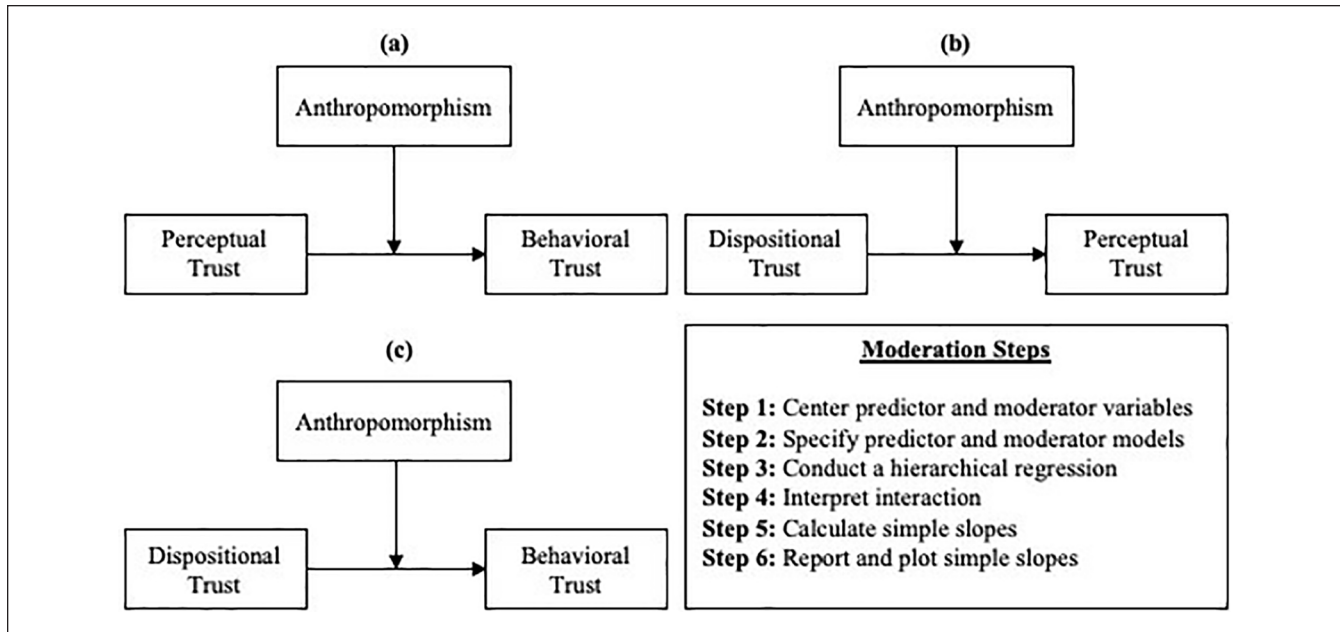


Figure 2. Moderation path diagrams and steps.

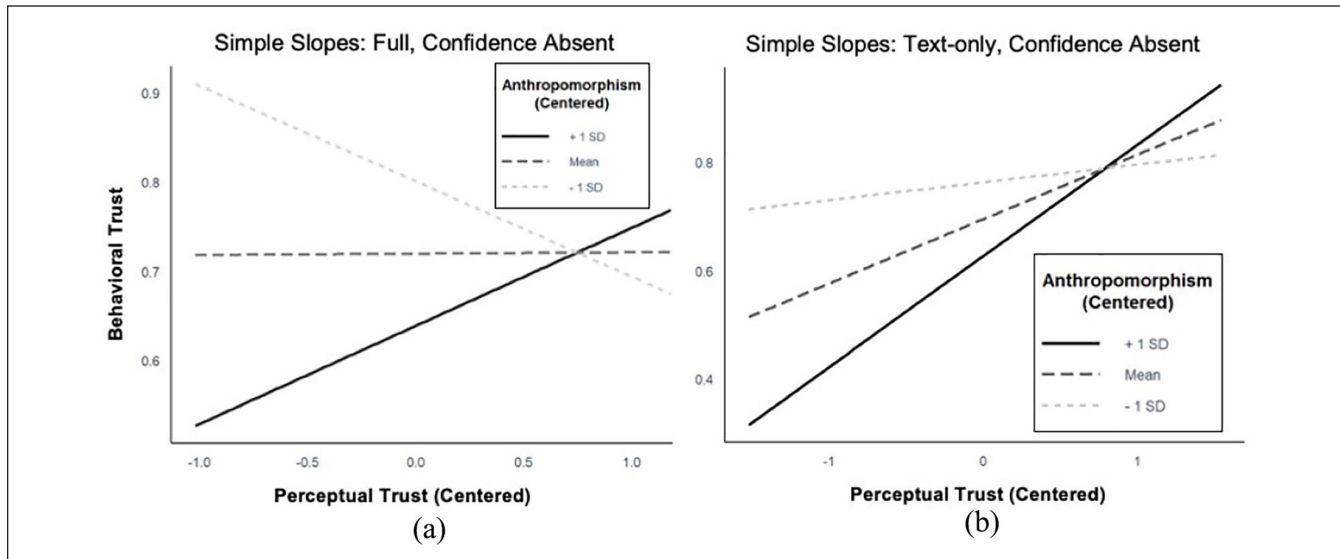


Figure 3. Simple slopes analysis of the relationship between perceptual and behavioral trust in the (a) Full, Confidence Absent condition, and in (b) the Textonly, Confidence Absent condition.

behavioral trust. A total of 18 moderation analyses—three for each condition—were conducted following J. Cohen et al. (2002), as summarized in Figure 2. Following similar team studies in STE settings (Cooke et al., 2007; Salem et al., 2013), a significance level of 0.10 was selected for this study.

Perceptual and Behavioral Trust

Hierarchical regressions for perceptual and behavioral trust resulted in significant moderations of anthropomorphism for

participants in the Full, Confidence Absent ($\Delta R^2 = 0.28$, $F(1, 7) = 9.52$, $p < .05$) and Text-only, Confidence Absent ($\Delta R^2 = 0.35$, $F(1, 5) = 14.84$, $p < .05$) conditions. No other significant moderations were found in other conditions.

Simple slopes analyses were conducted to elucidate the significant moderations (Figure 3). In the Full, Confidence Absent condition (Figure 3a), the significant moderator model and interaction term indicated that as anthropomorphism increases, the slope of the perceptual to behavioral trust relationship increases by $b = 0.10$, $t(7) = 3.09$, $p < .05$;

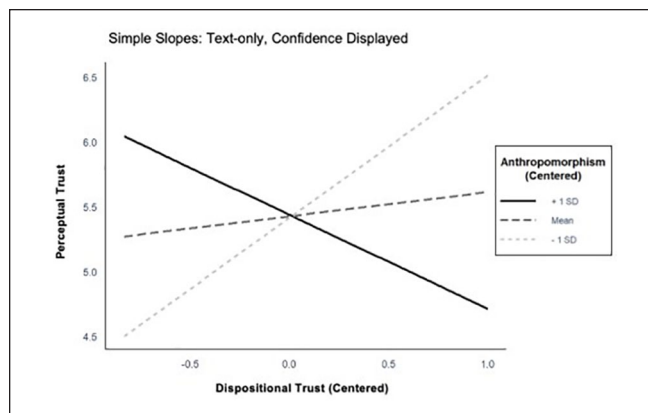


Figure 4. Simple slopes analysis of the relationship between dispositional and perceptual trust in the Text-only, Confidence Displayed condition.

$R^2 = 0.80$, $F(3, 7) = 9.16$, $p < .01$. Though the simple slope line at the mean level of anthropomorphism was not significant ($p = .946$), significance was found for those at +1SD with a 0.11 per-unit effect ($a = .64$, $b = .11$, $t(7) = 2.44$, $p < .05$), and at -1SD with a -0.11 per-unit effect ($a = .80$, $b = -.11$, $t(7) = -2.96$, $p < .05$).

In the Text-only, Confidence Absent condition (Figure 3b), the significant moderator model and interaction term indicated that as anthropomorphism increases, the slope increases by $b = 0.08$, $t(5) = 3.85$, $p < .05$; $R^2 = 0.88$, $F(3, 5) = 12.28$, $p = .01$. Simple slope analysis yielded significant slope lines at the mean level of anthropomorphism with a .12 per-unit increase effect ($a = 0.70$; $b = 0.12$, $t(5) = 4.78$, $p < .01$), and at +1SD with a .21 per-unit increase effect ($a = 0.63$; $b = 0.21$, $t(5) = 6.06$, $p < .01$).

Dispositional and Perceptual Trust

Anthropomorphism was a significant moderator of the relationship between dispositional and perceptual trust in the Text-only, Confidence Displayed condition (Figure 4), based on a significant hierarchical analysis and interaction term ($\Delta R^2 = 0.48$, $F(1, 9) = 27.84$, $p < .001$). No other significant moderations were found for other conditions. As anthropomorphism increases, the slope of the dispositional to perceptual trust relationship decreases by $b = -0.90$, $t(9) = -5.28$, $p < .001$; $R^2 = 0.84$, $F(3, 9) = 16.28$, $p < .001$. Simple slopes analyses revealed that although the simple slope at the mean level of anthropomorphism was not significant ($p = .232$), significance was found for slope lines for individuals at +1SD with a -0.73 per-unit effect ($a = 5.44$, $b = -0.73$, $t(9) = -2.77$, $p < .05$), and at -1SD with a 1.10 per-unit effect ($a = 5.41$, $b = 1.10$, $t(9) = 6.01$, $p < .001$).

Dispositional and Behavioral Trust

Anthropomorphism was found to be a significant moderator for individuals in the Text-only, Confidence Displayed,

$\Delta R^2 = 0.26$, $F(1, 9) = 4.68$, $p < .10$, and Text-only, Confidence Absent ($\Delta R^2 = 0.30$, $F(1, 5) = 4.07$, $p = .10$) conditions.

In the Text-only, Confidence Displayed condition, the significant moderator model and interaction term indicated that as perceived anthropomorphism increases, the slope decreases by $b = -0.11$, $t(9) = -2.16$, $p < .10$; $R^2 = 0.50$, $F(3, 9) = 3.02$, $p < .10$. The associated simple slopes analysis revealed a significant slope for individuals at +1 SD (Figure 5a), with a per-unit effect of -0.16 ($a = 0.78$; $b = -0.16$, $t(9) = -2.04$, $p < .10$).

In the Text-only, Confidence Absent condition indicated that as anthropomorphism increases, the slope increases by $b = 0.10$, $t(5) = 2.02$, $p = .10$. The moderator model which added the interaction was not significant ($R^2 = 0.64$, $F(3, 5) = 2.02$, $p = .139$). The associated simple slopes analysis revealed a significant slope for individuals at +1SD (Figure 5b) with a perunit effect of 0.19 ($a = 0.62$; $b = 0.19$, $t(5) = 2.96$, $p < .05$).

Discussion

Anthropomorphism significantly moderated the relationship between perceptual and behavioral trust in both the Full and Text-only conditions for Confidence Absent teams, partially supporting Hypothesis 1. In the absence of direct markers of the robot's confidence, more anthropomorphic perceptions related to the robot's lexical communication may have made participants more likely to rely on its perceived general trustworthiness as a heuristic for complying with each of its recommendations. This is consistent with previously reported "politeness" effects, in which people who interact with a machine more anthropomorphically tend to respond to its decisions more favorably (Nass, 2004). Surprisingly, our results also suggest that for participants who anthropomorphized less, more trustworthy perceptions resulted in lower compliance when recommendations were communicated textually. Thus, tempering anthropomorphism and using text-based communication may foster trustworthy perceptions about a robot without sacrificing the ability to scrutinize its immediate accuracy—even when it is unable to communicate its confidence in its decisions.

Anthropomorphism was also a significant moderator between dispositional and perceptual trust for participants in the Text-only, Confidence Displayed condition, partially supporting Hypothesis 2. It was also a significant moderator between dispositional and behavioral trust for Text-only teams in both Confidence manipulations, partially supporting Hypothesis 3. For participants who anthropomorphized the robot more as it textually communicated recommendations with confidence indicators, dispositional trust was inversely proportional to both perceptual and behavioral trust. Communicating uncertainty is a social etiquette that people expect from machines during critical interactions (Parasuraman & Miller, 2004). Participants who were anthropomorphizing the robot more may have interpreted its

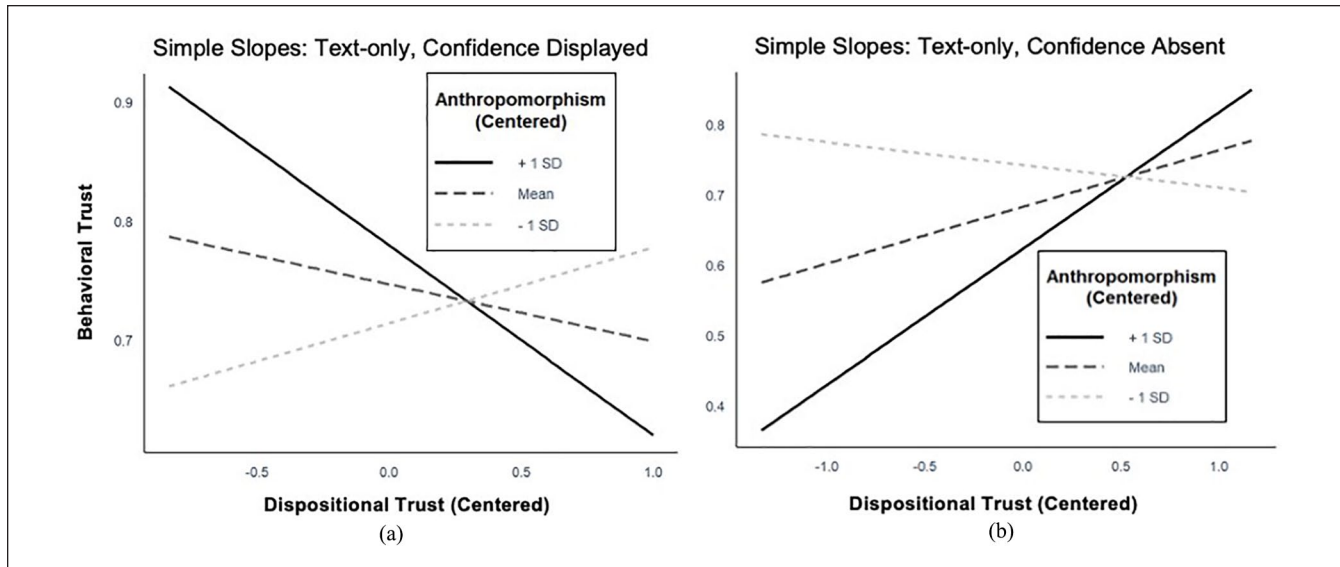


Figure 5. Simple slopes analysis of the relationship between dispositional and behavioral trust in the (a) Text-only, Confidence Displayed condition, and (b) the Text-only, Confidence Absent condition.

expressions of uncertainty to mean that trusting it less is an appropriate response, even if they tended to generally trust robots. Curiously, dispositional trust was directly proportional to perceptual trust for low-anthropomorphism participants when confidence was displayed, and to behavioral trust for high-anthropomorphism participants when not. Thus, to build trust based on a robot's contextual reliability instead of a person's general opinions about robots, confidence communication may have to be accompanied by a minimum level of anthropomorphic perceptions.

We note that all hypothesized moderations were found to be significant only for participants in conditions that involved the robot communicating by text. In our study, lexical communication in the Text-only and Full conditions may have aroused sufficiently strong beliefs about the human-likeness of the robot's cognitive abilities to affect how dispositional trust translated to perceptual trust, and subsequently, to behavioral trust. Therefore, language-based communication might serve as a conduit for anthropomorphic perceptions to influence the relationship between various forms of trust. Further work should investigate if similar moderations occur when information is communicated lexically in other modalities, such as voice.

We acknowledge some limitations of this study. First, the limited sample sizes per group resulted in non-normal data and low power, potentially limiting the generalizability of our findings. The administration of the Godspeed scales as semantic differential scales instead of Likert scales (as in Kaplan et al., 2021) may have also amplified the role of individual differences in anthropomorphism responses. Finally, our remote data collection in Roblox might have accentuated

dispositional trends from participants who tended to be younger and perhaps less likely to anthropomorphize robots (Letheren et al., 2016).

Conclusion

This study demonstrated that anthropomorphism moderates the relationships between dispositional, perceptual, and behavioral trust in a virtual robot teammate. Our findings suggest complex relationships between different conceptualizations of trust, anthropomorphism, and robot communication styles, including the communication of confidence information. The role of anthropomorphic perceptions should, therefore, be considered in designing for and evaluating how language-based robot communication features affect trust. Interactive team cognition theory (Cooke et al., 2013) suggests that factors like trust and anthropomorphism may evolve dynamically as teammates observe team cognitive artifacts that arise from their interactions. Thus, studies on teams with more than two members may benefit from exploring trust and anthropomorphism through interactive communication measures, such as the usage of using personifying or objectifying references to a robot (M. C. Cohen et al., 2021). Finally, future research should consider how language-based confidence communication affects HRT processes and performance as anthropomorphism moderates the relationships between dispositional, perceptual, and behavioral trust.

Acknowledgments

This research was partially supported by a grant from the Air Force Office of Scientific Research [FA9550-18-1-0067].

ORCID iD

Myke C. Cohen  <https://orcid.org/0000-0003-0520-4507>

References

- Bartneck, C., Kulić, D., Croft, E., & Zoghbi, S. (2009). Measurement Instruments for the Anthropomorphism, Animacy, Likeability, Perceived Intelligence, and Perceived Safety of Robots. *International Journal of Social Robotics*, 1(1), 71–81. <https://doi.org/10.1007/s12369-008-0001-3>
- Casper, J., & Murphy, R. R. (2003). Human-robot interactions during the robot-assisted urban search and rescue response at the World Trade Center. *IEEE Transactions on Systems, Man, and Cybernetics, Part B (Cybernetics)*, 33(3), 367–385. <https://doi.org/10.1109/TSMCB.2003.811794>
- Chancey, E. T., Bliss, J. P., Yamani, Y., & Handley, H. A. H. (2017). Trust and the Compliance–Reliance Paradigm: The Effects of Risk, Error Bias, and Reliability on Trust and Dependence. *Human Factors*, 59(3), 333–345. <https://doi.org/10.1177/0018720816682648>
- Chiou, E. K., & Lee, J. D. (2023). Trusting Automation: Designing for Responsivity and Resilience. *Human Factors: The Journal of the Human Factors and Ergonomics Society*, 65(1), 137–165. <https://doi.org/10.1177/00187208211009995>
- Cohen, J., Cohen, P., West, S. G., & Aiken, L. S. (2002). *Applied Multiple Regression/Correlation Analysis for the Behavioral Sciences* (3rd ed.). Routledge. <https://doi.org/10.4324/9780203774441>
- Cohen, M. C., Demir, M., Chiou, E. K., & Cooke, N. J. (2021). The Dynamics of Trust and Verbal Anthropomorphism in Human-Autonomy Teaming. *2021 IEEE 2nd International Conference on Human-Machine Systems (ICHMS)*, 1–6. <https://doi.org/10.1109/ICHMS53169.2021.9582655>
- Cooke, N. J., Cohen, M. C., Fazio, W. C., Inderberg, L. H., Johnson, C. J., Lematta, G. J., Peel, M., & Teo, A. (2023). From Teams to Teamness: Future Directions in the Science of Team Cognition. *Human Factors*. <https://doi.org/10.1177/00187208231162449>
- Cooke, N. J., Gorman, J. C., Duran, J. L., & Taylor, A. R. (2007). Team cognition in experienced command-and-control teams. *Journal of Experimental Psychology: Applied*, 13(3), 146–157. <https://doi.org/10.1037/1076-898X.13.3.146>
- Cooke, N. J., Gorman, J. C., Myers, C. W., & Duran, J. L. (2013). Interactive Team Cognition. *Cognitive Science*, 37(2), 255–285. <https://doi.org/10.1111/cogs.12009>
- Cooke, N. J., & Shope, S. M. (2004). Designing a Synthetic Task Environment. In S. G. Schifflett, L. R. Elliott, E. Salas, & M. D. Covert (Eds.), *Scaled Worlds: Development, Validation and Applications* (pp. 263–278).
- de Visser, E. J., Monfort, S. S., McKendrick, R., Smith, M. A. B., McKnight, P. E., Krueger, F., & Parasuraman, R. (2016). Almost human: Anthropomorphism increases trust resilience in cognitive agents. *Journal of Experimental Psychology: Applied*, 22(3), 331–349. <https://doi.org/10.1037/xap0000092>
- de Visser, E. J., Peeters, M. M. M., Jung, M. F., Kohn, S., Shaw, T. H., Pak, R., & Neerincx, M. A. (2020). Towards a Theory of Longitudinal Trust Calibration in Human–Robot Teams. *International Journal of Social Robotics*, 12(2), 459–478. <https://doi.org/10.1007/s12369-019-00596-x>
- Epley, N., Waytz, A., & Cacioppo, J. T. (2007). On Seeing Human: A Three-Factor Theory of Anthropomorphism. *Psychological Review*, 114(4), 864–886. <https://doi.org/10.1037/0033-295x.114.4.864>
- Fischer, K. (2011). Interpersonal variation in understanding robots as social actors. *2011 6th ACM/IEEE International Conference on Human-Robot Interaction (HRI)*, 53–60. <https://doi.org/10.1145/1957656.1957672>
- Groom, V., & Nass, C. (2007). Can robots be teammates?: Benchmarks in human–robot teams. *Interaction Studies*, 8(3), 483–500. <https://doi.org/10.1075/is.8.3.10gro>
- Hancock, P. A., Billings, D. R., Schaefer, K. E., Chen, J. Y. C., de Visser, E. J., & Parasuraman, R. (2011). A Meta-Analysis of Factors Affecting Trust in Human-Robot Interaction. *Human Factors: The Journal of the Human Factors and Ergonomics Society*, 53(5), 517–527. <https://doi.org/10.1177/0018720811417254>
- Jensen, T., Khan, M. M. H., & Albayram, Y. (2020). The Role of Behavioral Anthropomorphism in Human-Automation Trust Calibration. In H. Degen & L. Reinerman-Jones (Eds.), *Artificial Intelligence in HCI* (pp. 33–53). Springer International Publishing. https://doi.org/10.1007/978-3-030-50334-5_3
- Jessup, S. A., Schneider, T. R., Alarcon, G. M., Ryan, T. J., & Capiola, A. (2019). The Measurement of the Propensity to Trust Automation. In J. Y. C. Chen & G. Fragomeni (Eds.), *Virtual, Augmented and Mixed Reality. Applications and Case Studies* (pp. 476–489). Springer International Publishing. https://doi.org/10.1007/978-3-030-21565-1_32
- Kaplan, A. D., Sanders, T. L., & Hancock, P. A. (2021). Likert or Not? How Using Likert Rather Than Bipolar Ratings Reveal Individual Difference Scores Using the Godspeed Scales. *International Journal of Social Robotics*, 13(7), 1553–1562. <https://doi.org/10.1007/s12369-020-00740-y>
- Kohn, S. C., de Visser, E. J., Wiese, E., Lee, Y.-C., & Shaw, T. H. (2021). Measurement of Trust in Automation: A Narrative Review and Reference Guide. *Frontiers in Psychology*, 12. <https://www.frontiersin.org/articles/10.3389/fpsyg.2021.604977>
- Kulms, P., & Kopp, S. (2019). More Human-Likeness, More Trust?: The Effect of Anthropomorphism on Self-Reported and Behavioral Trust in Continued and Interdependent Human-Agent Cooperation. *Proceedings of Mensch Und Computer 2019*, 31–42. <https://doi.org/10.1145/3340764.3340793>
- Lee, J. D., & See, K. A. (2004). Trust in Automation: Designing for Appropriate Reliance. *Human Factors*, 46(1), 50–80. https://doi.org/10.1518/hfes.46.1.50_30392
- Letheren, K., Kuhn, K.-A. L., Lings, I., & Pope, N. K. Ll. (2016). Individual difference factors related to anthropomorphic tendency. *European Journal of Marketing*, 50(5/6), 973–1002. <https://doi.org/10.1108/EJM-05-2014-0291>
- Mayer, R. C., Davis, J. H., & Schoorman, F. D. (1995). An Integrative Model Of Organizational Trust. *Academy of Management Review*, 20(3), 709–734. <https://doi.org/10.5465/amr.1995.9508080335>
- Meyer, J., & Lee, J. D. (2013). Trust, reliance, and compliance. In J. D. Lee & A. Kirlik (Eds.), *The Oxford handbook of cognitive engineering* (pp. 109–124). Oxford University Press.

- Mori, M. (1970). Bukimi no tani (the uncanny valley). *Energy*, 7(4), 33–35.
- Nass, C. (2004). Etiquette equality: Exhibitions and expectations of computer politeness. *Communications of the ACM*, 47(4), 35–37. <https://doi.org/10.1145/975817.975841>
- Nass, C., & Moon, Y. (2000). Machines and Mindlessness: Social Responses to Computers. *Journal of Social Issues*, 56(1), 81–103. <https://doi.org/10.1111/0022-4537.00153>
- Parasuraman, R., & Miller, C. A. (2004). Trust and etiquette in high-criticality automated systems. *Communications of the ACM*, 47(4), 51–55. <https://doi.org/10.1145/975817.975844>
- Raimondo, F. R., Wolff, A. T., Hehr, A. J., Peel, M. A., Wong, M. E., Chiou, E. K., Demir, M., & Cooke, N. J. (2022). Trailblazing Roblox Virtual Synthetic Testbed Development for Human-Robot Teaming Studies. *Proceedings of the Human Factors and Ergonomics Society Annual Meeting*, 66, 812–816. <https://doi.org/10.1177/10711813222661470>
- Riek, L. (2012). Wizard of Oz Studies in HRI: A Systematic Review and New Reporting Guidelines. *Journal of Human-Robot Interaction*, 119–136. <https://doi.org/10.5898/JHRI.1.1.Riek>
- Salem, M., Eyssel, F., Rohlfing, K., Kopp, S., & Joublin, F. (2013). To Err is Human(-like): Effects of Robot Gesture on Perceived Anthropomorphism and Likability. *International Journal of Social Robotics*, 5(3), 313–323. <https://doi.org/10.1007/s12369-013-0196-9>
- Stein, J.-P., & Ohler, P. (2017). Venturing into the uncanny valley of mind— The influence of mind attribution on the acceptance of human-like characters in a virtual reality setting. *Cognition*, 160, 43–50. <https://doi.org/10.1016/j.cognition.2016.12.010>
- Waytz, A., Heafner, J., & Epley, N. (2014). The mind in the machine: Anthropomorphism increases trust in an autonomous vehicle. *Journal of Experimental Social Psychology*, 52, 113–117. <https://doi.org/10.1016/j.jesp.2014.01.005>