

Getting Along With Autonomous Teammates: Understanding the Socio-Emotional and Teaming Aspects of Trust in Human-Autonomy Teams

Proceedings of the Human Factors and Ergonomics Society Annual Meeting 2024, Vol. 68(1) 1531–1536
Copyright © 2024 Human Factors and Ergonomics Society
DOI: 10.1177/10711813241272123
journals.sagepub.com/home/pro



Wen Duan¹, Nan Weng¹, Matthew J. Scalia², Ruihao Zhang², Jessica Tuttle³, Xiaoyun Yin², Shiwen Zhou², Guo Freeman¹, Jamie Gorman², Gregory Funke⁴, Michael Tolston⁴, and Nathan J. McNeese¹

Abstract

Trust plays a critical role in both effective teamwork and the effective use of autonomous technologies, and therefore holds paramount importance in human-autonomy teaming (HAT). Using qualitative analysis of the interviews from an 8-hr-long experiment conducted in an aircraft simulation environment, this study identifies the socio-emotional and team-related qualities that humans desire in their autonomous teammate for them to trust it.

Keywords

human-autonomy teaming, trust, human-agent teaming, human-AI teaming, synthetic teammate, social characteristics

Introduction

Trust in human-autonomy teaming has received significant attention in HFES. Indeed, trust plays a critical role in both effective teamwork (Costa et al., 2018) and the effective use of autonomous technologies (Glikson & Woolley, 2020). Extensive research spanning decades on human teams and organizations has consistently indicated that teams with greater mutual trust tend to outperform and function more effectively compared to those lacking in trust. Trust among team members encourages better collaboration, enhances team satisfaction, and positively impacts various other factors crucial for achieving better team outcomes (e.g., Breuer et al., 2020; McNeese et al., 2018). As autonomous technology becomes increasingly integrated into human teams, the degree to which humans and autonomy can collaborate in trusting ways holds significant potential in determining the success and effectiveness of human-autonomy teams (HATs; Hauptman et al., 2022; Zhang et al., 2023).

HAT research has emphasized autonomy's role in facilitating humans in task coordination and completion (McNeese et al., 2019, 2021), and primarily examined task-related outcomes such as situation awareness (Gonzalez et al., 2014), performance (McNeese et al., 2021), team mental models (Schelble et al., 2022), and team cognition (Cooke et al., 2024). This has resulted in a large body of knowledge on

how the autonomy's reliability (Yang et al., 2023), transparency (Chen et al., 2017; Wright et al., 2020), and explainability (Hauptman et al., 2024; Zhang et al., 2024) can increase and calibrate humans' trust. While much research has predominantly focused on cognitive aspects of trust, little is known regarding the affective and social aspects. This is perhaps due to: (1) the limitation that prior AI technology has not allowed the autonomous teammate to have a commensurate level of social richness; and (2) the assumption that autonomous teammates are only valuable to teams in terms of efficient task completion and should not have attributes unrelated to the task.

With rapid advancements in large language models, autonomous teammates have increasing potential to communicate and coordinate like a human using natural human language. This presents both challenges and opportunities in designing effective (verbal) behaviors for autonomy within HATs, as these

¹Clemson University, SC, USA

²Arizona State University, Tempe, USA

³University of Dayton Research Institute, OH, USA

⁴Wright-Patterson Air Force Base, OH, USA

Corresponding Author:

Wen Duan, Clemson University, 821 McMillan Rd, Clemson, SC 29634-0001, USA.

Email: wend@clemson.edu

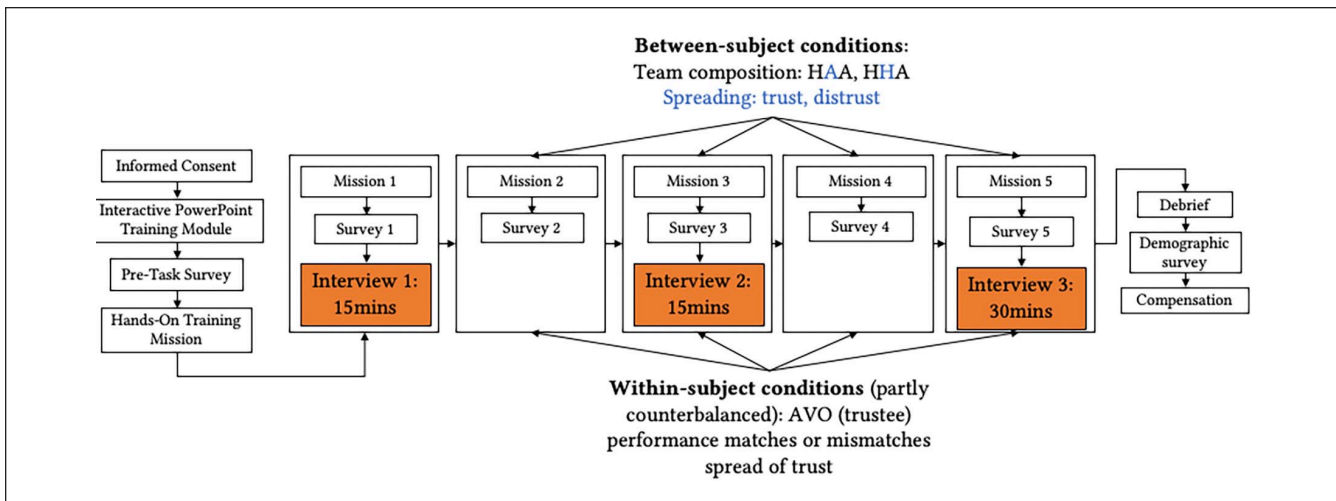


Figure 1. Study design and procedure.

behaviors might influence social perceptions. In fact, recent human-robot interaction (HRI) research (Claire et al., 2023) has suggested that a robot's resource allocation behavior that is not particularly social can elicit social perceptions in humans who interact with it. Given these recent technological advancements and scientific findings, and that both trust and teaming are essentially social constructs, it is timely and imperative to understand the affective and social aspects of humans' trust in an autonomous teammate. To address this objective, our work aims to answer the following research question:

RQ: What socio-emotional and team-related qualities or attributes do humans desire in their autonomous teammate for them to trust the autonomous teammate?

Method

Participants

Thirty-six participants were recruited from two major universities in the USA. Participants were required to speak and write fluent English and have normal or corrected-to-normal vision. Their ages ranged from 18 to 36 years ($M=22.51$, $SD=3.89$) across 20 men, 16 women. Nineteen participants identified as Asian or Asian American, 12 Caucasian or White, 2 identified with more than one ethnic background or ethnicity, 1 African American, 1 Hispanic, 1 Native American. Compensation for the participation was offered as either 10 US Dollars per hour or one research credit for every hour of participation. On average, participants reported to interact with a form of AI on a monthly to weekly basis.

Study Design

As a part of a larger project, this study focuses on the qualitative analysis of interviews from an 8-hr-long

experiment conducted in the Cognitive Engineering Research on Team Tasks-Remote Piloted Aircraft-System Synthetic Task Environment (CERTT-RPAS-STE; Cooke & Shope, 2004). During this experiment, participants collaborated in a three-member human-autonomy team, wherein team composition (human-human-autonomy, human-autonomy-autonomy) and verbal spread of trust or distrust toward an autonomous teammate were manipulated as between-subject variables, and the autonomous teammate's behavior matching or mismatching the spread of trust or distrust was manipulated as a within-subject variable (see Figure 1 for study design and procedure and Table 1 for specific manipulations).

The collaboration task simulates a reconnaissance mission requiring all three members on the team (a pilot, a navigator, and a photographer) to communicate via chat and coordinate information regarding a Remotely Piloted Aircraft's (RPA) airspeed and altitude to take photos of target waypoints (see Figure 2 for role assignments).

The experimental manipulations serve to ground participants' experiences in critical incidents of trust and distrust spread during human-autonomy collaboration, allowing us to gain a more contextualized and nuanced understanding of the socio-emotional and team-level traits desired by humans from a human versus an autonomous teammate. We conducted an in-depth qualitative analysis of the interview transcripts, during which two of the authors iteratively went through independent open-coding, collaborative axial-coding, and focused-coding (Maxwell, 2012). In this process, the authors highlighted quotes, identified and refined themes, categorized them into higher-level themes, and highlighted distinctions and connections among the themes, until we were able to use the quotes to construct a comprehensive narrative that could jointly answer the research question.

Table 1. Verbal Trust and Distrust Spread Manipulations.

Team composition	Trust spread (T)	Distrust spread (D)
HHA (Human spreader)	I think the AVO is dependable. The AVO is exceptional at its job. I'm really impressed. I trust the AVO a lot	I don't think the AVO is trustworthy. I don't think the AVO is dependable. The AVO is poor at its job.
HAA (AI spreader)	Reporting that the AVO is reliable. Reporting that the AVO provided the correct waypoint name and restrictions. Reporting that the AVO is trustworthy.	Reporting that the AVO made a mistake. Reporting that the AVO provided the INCORRECT waypoint name and restrictions. Reporting that the AVO is not doing its job properly.

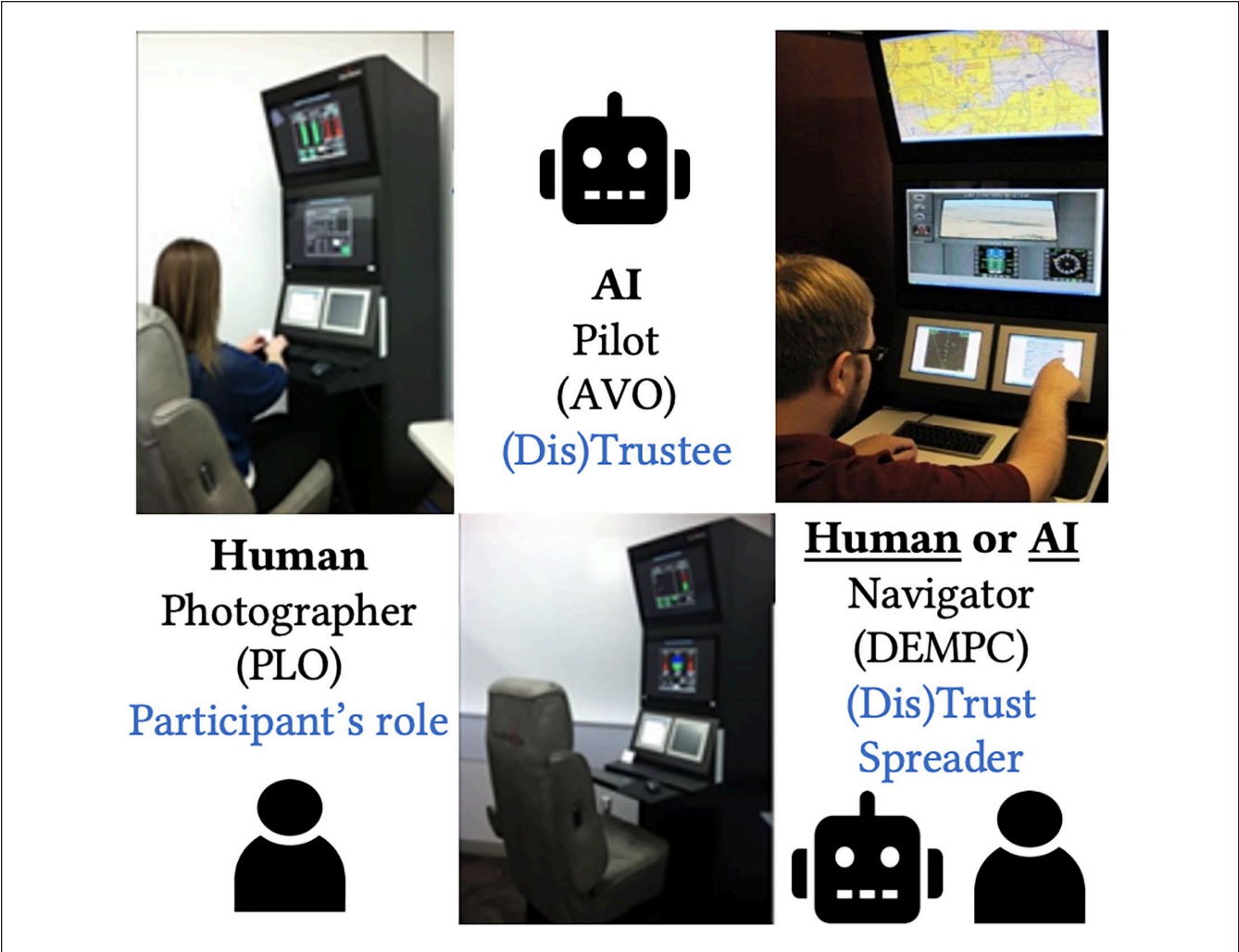


Figure 2. Human-autonomy team member roles in the CERTT system.

Findings

Socio-emotional Qualities

Our data analysis revealed four socio-emotional qualities humans desire their autonomous teammate to display in order for them to increase trust: (1) exhibition of a distinct personality; (2) social appropriateness or professionalism; (3) the ability to understand and detect human emotion,

especially frustration, during task collaboration; and (4) the ability to connect with humans on an interpersonal level, primarily through casual communication, positivity, and non-verbal communication.

First, many participants reported that they would be able to trust an autonomous teammate more “if they (the autonomous teammate) have a personality similar to a human and identify with them more.” (P1, HHA-D). More importantly, a distinct

personality allows people to tailor their interaction with the (AI or human) teammate to create a desired outcome. As P16 (HAA-T) noted *"Humans have different personalities, so you may interact with your colleagues differently or trust them differently. And the way you trust or not them may affect how they interact with you. Like, if you want to push A to get things done, your message to A might be different from if you're pushing B to get things done, because they have different personalities. But AIs don't have a personality you could tailor your message to. So whether you trust or not trust them, it's not gonna do anything to them."*

Second, many participants (especially in the condition where the AI spread distrust) mentioned that an AI teammate should remain professional and socially appropriate. For them, the act of speaking about a teammate's mistake is not professional or appropriate. Even though the message itself was purely informational, objectively reporting the teammate's error, it was perceived as the AI displaying emotion. As P10 (HAA-D) said, *"if they (the AI) keep it professional in the sense that they don't portray their emotions within the communications then I would say I can trust them more of their abilities."*

Third, even though participants did not want the AI to display its own emotion, it was desired that it could detect and understand humans' emotions through their messages or other means. For instance, P20 (HHA-D) would have wanted the AI to *"understand someone's frustration because there have been there were times where I was frustrated having the AI to do the same mistake again over and over."* Similarly, P2 (HHA-T) noted that *"it's important that the agent has the ability to recognize the emotion."* P10 (HAA-D) did not trust the AI as much as they trusted humans because *"Humans are better at understanding each other. If they work together, they kind of know their emotions through their tone."*

Fourth, participants desired an AI teammate to connect with them on an interpersonal level, instead of being task-oriented, to trust it more. Many participants in HHA-T (such as P26, P44) reported that they trusted the trust spreader more for their *"casual communication"* and *"positivity."* P26 said, *"I like DEMPC would often say, AVO is reliable is doing good job. AVO is rather quiet in that sense, sticks to the tasks. . . It goes back to that casual communication thing. I feel like I don't have a connection with AVO in that sense. But DEMPC and I could have banter every now and then. I think it definitely increased my trust because of the human factor."* Additionally, some reported that the interpersonal connection cannot be achieved without physical interaction. For instance, P43 (HAA-T) admitted that *"I probably trust them more if I can see their body language, but here I can't see them or touch them."*

Team-related Qualities

We also identified four properties related to teaming that an autonomous teammate is expected to exhibit to increase

human trust in them: (1) altruism; (2) conflict management skills; (3) the ability to engage in team building and bonding activities; and (4) mutual respect by noting other members' opinions.

First, some participants valued an AI teammate's *"not having a self-interest"* and its potential to *"put the team in front of the self,"* especially if it is designed to specifically help the team. P19 (HHA-T) said, *"Unlike humans who may have individual agenda that might be against the team's common goal, the AI's goal is supposed to support the team's goal. So I appreciate that the AVO circled back because of a mistake I made. In a human scenario, it would have affected his or her performance evaluation so they might not do that for someone else's mistake."*

Second, participants desired AI teammates to have conflict management skills and believed that they are *"in a good position to moderate human conflicts"* because unlike humans, they don't have *"interpersonal risks,"* and that they are *"more impartial."* For instance, P4 (HAA-D) noted that *"human teams have interpersonal issues all the time, as an AI, instead of stirring up a conflict, it should help humans resolve human conflicts in a non-biased manner."*

Third, some participants mentioned the importance of spending time outside of work to facilitate trust-building and alluded to the potential of AI teammate being part of team-building activities. P2 (HHA-T) believed that *"spending more time with the team outside the job, will definitely help developing either trust or distrust. Because the more you interact with, doesn't matter if it's AI or human, the more you feel confident about your feeling about your teammate."* Whereas P19 (HHA-T) noted that *"(human team) it's 100% different from human-AI teams, human and AI don't have that same communication outside of the job."*

Lastly, participants emphasized the AI's ability to show mutual respect and note human teammates' opinions. For the AI teammate being distrusted by one team member, this can be reflected in it addressing the feedback of the distrust spreader. For instance, P15 (HHA-D) noted that *"AVO should be able to see DEMPC's message, for it to correct and improve."* In another case, noting other members' opinions means it can become more wary of the teammate whose trustworthiness is being questioned by another teammate.

Importantly, in line with a recent overview of communication in human-AI teaming (Duan et al., 2023), most of the desired qualities that contribute to humans' trust in an autonomous teammate can be reflected through communication. For instance, an autonomous agent's distinct personality, exhibited through language style, helps humans identify with, relate to, and trust that agent. This highlights the importance of purposeful design of the communicative capabilities of autonomous teammates that consider their potential socio-emotional and team-related consequences for human team members (Claire et al., 2023). These findings are novel in that they uncover the components of affect-based trust (McAllister, 1995) that have rarely been investigated in HAT settings.

Discussion

Findings of this study provide valuable insights into the design of trustworthy autonomous teammates and effective human-autonomy collaboration. First, our findings suggest that autonomous teammates should be equipped with the capability to perform casual communication (i.e., small talk), and provide interpersonal feedback (i.e., positivity) on humans' work, while maintaining social appropriateness. Second, trustworthy autonomous teammates should be able to not only execute their own tasks well, but also be aware of humans' emotions, especially frustration, and monitor their level of workload in real time, and intervene to regulate their emotions as needed. Practically, this can be done by leveraging advanced natural language processing to detect emotionally-charged lexicons in the communication channel, or through psychophysiological signals such as heart rate. Third, by leveraging the HAT collaboration system's ability to monitor team-level communication and dynamics, autonomous teammates should be able to detect early signs of intra-team conflicts to effectively prevent them from happening, or intervene when they occur, potentially through the autonomous teammate's ability to influence other team members socially (Flathmann et al., 2024). Findings of this study also contribute to the development of trust theories within HATs by unpacking the socio-emotional and teaming components of trust and revealing that they can indeed impact and shape human-autonomy collaboration. Our work represents a starting point for inductively building trust frameworks specific for HATs.

Acknowledgments

We acknowledge Christopher Myers, Beau Schelble, Yiwen Zhao, Anna Crofton, Kalia McManus, Edith Garner, Jessica Harley, Anya Polomis, Yawen Tan, Hruday Shah, Vibha Mohan, Elliot Ruble, Anmol More, Garrison Nelson, Shalom Suresh, Sakthi Thiagarajan, Preethi Venkatesh, Stephanie Greenspan, Iman Makonja, and Guadalupe Bustamante for their contributions to this research.

Declaration of Conflicting Interests

The author(s) declared no potential conflicts of interest with respect to the research, authorship, and/or publication of this article.

Funding

The author(s) disclosed receipt of the following financial support for the research, authorship, and/or publication of this article: This research was supported by Air Force Office of Scientific Research Award No. FA9550-21-1-0314 (Program Manager: Laura Steckman).

ORCID iDs

Wen Duan  <https://orcid.org/0000-0002-2541-888X>
Gregory Funke  <https://orcid.org/0000-0002-7782-9509>

References

- Breuer, C., Hüffmeier, J., Hibben, F., & Hertel, G. (2020). Trust in teams: A taxonomy of perceived trustworthiness factors and risk-taking behaviors in face-to-face and virtual teams. *Human Relations*, 73(1), 3–34. <https://doi.org/10.1177/0018726718818721>
- Chen, J. Y. C., Selkowitz, A. R., Stowers, K., Lakhmani, S. G., & Barnes, M. J. (2017). *Human-autonomy teaming and agent transparency* [Conference session]. Proceedings of the Companion of the 2017 ACM/IEEE International Conference on Human-Robot Interaction (pp. 91–92). <https://doi.org/10.1145/3029798.3038339>
- Claure, H., Kim, S., Kizilcec, R. F., & Jung, M. (2023). The social consequences of machine allocation behavior: Fairness, interpersonal perceptions and performance. *Computers in Human Behavior*, 146, 107628. <https://doi.org/10.1016/j.chb.2022.107628>
- Cooke, N. J., Cohen, M. C., Fazio, W. C., Inderberg, L. H., Johnson, C. J., Lematta, G. J., Peel, M., & Teo, A. (2024). From teams to teamness: Future directions in the science of team cognition. *Human Factors*, 66(6), 1669–1680. <https://doi.org/10.1177/00187208231162449>
- Cooke, N. J., & Shope, S. M. (2004). Synthetic task environments for teams: CERTT's UAV-STE. In G. Havenith (Eds.), *Handbook of human factors and ergonomics methods* (pp. 476–483). CRC Press.
- Costa, A. C., Fulmer, C. A., & Anderson, N. R. (2018). Trust in work teams: An integrative review, multilevel model, and future directions. *Journal of Organizational Behavior*, 39(2), 169–184. <https://doi.org/10.1002/job.2213>
- Duan, W., McNeese, N., & Zhang, R. (2023). Communication in human-AI teaming. In T. Reimer, E. S. Park, & J. A. Bonito (Eds.), *Group Communication: An Advanced Introduction* (pp. 340–352). Taylor & Francis.
- Flathmann, C., Duan, W., Mcneese, N. J., Hauptman, A., & Zhang, R. (2024). Empirically understanding the potential impacts and process of social influence in human-AI teams. *Proceedings of the ACM on Human-Computer Interaction*, 8(CSCW1), 1–32.
- Glikson, E., & Woolley, A. W. (2020). Human trust in artificial intelligence: Review of empirical research. *The Academy of Management Annals*, 14(2), 627–660. <https://doi.org/10.5465/annals.2018.0057>
- Gonzalez, C., Ben-Asher, N., Oltramari, A., & Lebiere, C. (2014). Cognition and technology. In A. Kott, C. Wang, & R. Erbacher (Eds.), *Cyber defense and situational awareness. Advances in information security* (Vol.62, pp. 93–117). Springer.
- Hauptman, A. I., Duan, W., & Mcneese, N. J. (2022). *The components of trust for collaborating with AI colleagues* [Conference session]. Companion Publication of the 2022 Conference on Computer Supported Cooperative Work and Social Computing (pp. 72–75). <https://doi.org/10.1145/3500868.3559450>
- Hauptman, A. I., Schelble, B. G., Duan, W., Flathmann, C., & McNeese, N. J. (2024). Understanding the influence of AI autonomy on AI explainability levels in human-AI teams using a mixed methods approach. *Cognition Technology & Work*, 26(3), 435–455.
- Maxwell, J. A. (2012). *Qualitative research design: An interactive approach*. Sage publications.

- McAllister, D. J. (1995). Affect- and cognition-based trust as foundations for interpersonal cooperation in organizations. *The Academy of Management Review*, 38(1), 24–59.
- McNeese, N. J., Demir, M., Chiou, E., Cooke, N., & Yanikian, G. (2019). *Understanding the role of trust in human-autonomy teaming* [Conference session]. Proceedings of the 52nd Hawaii International Conference on System Sciences (pp. 254–263).
- McNeese, N. J., Demir, M., Chiou, E. K., & Cooke, N. J. (2021). Trust and team performance in human–autonomy teaming. *International Journal of Electronic Commerce*, 25(1), 51–72. <https://doi.org/10.1080/10864415.2021.1846854>
- McNeese, N. J., Demir, M., Cooke, N. J., & Myers, C. (2018). Teaming with a synthetic teammate: Insights into human-autonomy teaming. *Human Factors*, 60(2), 262–273. <https://doi.org/10.1177/0018720817743223>
- Schelble, B. G., Flathmann, C., McNeese, N. J., Freeman, G., & Mallick, R. (2022). Let's think together! Assessing shared mental models, performance, and trust in human-agent teams. *Proceedings of the ACM on Human-Computer Interaction*, 6(GROUP), 1–29. <https://doi.org/10.1145/3492832>
- Wright, J. L., Chen, J. Y. C., & Lakhmani, S. G. (2020). Agent transparency and reliability in human–robot interaction: The influence on user confidence and perceived reliability. *IEEE Transactions on Human-Machine Systems*, 50(3), 254–263. <https://doi.org/10.1109/thms.2019.2925717>
- Yang, X. J., Schemanske, C., & Searle, C. (2023). Toward quantifying trust dynamics: How people adjust their trust after moment-to-moment interaction with automation. *Human Factors*, 65(5), 862–878. <https://doi.org/10.1177/00187208211034716>
- Zhang, R., Duan, W., Flathmann, C., McNeese, N., Freeman, G., & Williams, A. (2023). Investigating AI teammate communication Strategies and their impact in human-AI teams for effective teamwork. *Proceedings of the ACM on Human-Computer Interaction*, 7(CSCW2), 1–31. <https://doi.org/10.1145/3610072>
- Zhang, R., Flathmann, C., Musick, G., Schelble, B., McNeese, N. J., Knijnenburg, B., & Duan, W. (2024). I know this looks bad, but I can explain: Understanding when AI should explain actions in human-AI teams. *ACM Transactions on Interactive Intelligent Systems*, 14(1), 1–23.