

Assessing Communication and Trust in an AI Teammate in a Dynamic Task Environment

Shawaiz Bhatti

Human Systems Engineering
Arizona State University
Mesa, USA

sabhatti1@asu.edu

Mustafa Demir

Human Systems Engineering
Arizona State University
Mesa, USA

<https://orcid.org/0000-0002-5667-3701>

Nancy J. Cooke

Human Systems Engineering
Arizona State University
Mesa, USA

<https://orcid.org/0000-0003-0408-5796>

Craig J. Johnson

Human Systems Engineering
Arizona State University
Mesa, USA

cjohn42@asu.edu

Abstract— This research examines the relationship between anticipatory pushing of information and trust in human-autonomy teaming in a remotely piloted aircraft system - synthetic task environment. Two participants and one AI teammate emulated by a confederate executed a series of missions under routine and degraded conditions. We addressed the following questions: (1) How do anticipatory pushing of information and trust change from human to human and human to autonomous team members across the two sessions? and (2) How is anticipatory pushing of information associated with the trust placed in a teammate across the two sessions? This study demonstrated two main findings: (1) anticipatory pushing of information and trust differed between human-human and human-AI dyads, and (2) anticipatory pushing of information and trust scores increased among human-human dyads under degraded conditions but decreased in human-AI dyads.

Keywords—Artificial Intelligence, Communication, Human-Machine Teaming, Trust, Teamwork

I. INTRODUCTION

Technology continues to advance at a rapid rate. Much of this advancement can be seen in the field of machine learning, and artificial intelligence (AI), as highly autonomous machines (e.g., AI, robots, synthetic agents) permeate nearly all aspects of everyday life. There are many current examples of how AI and robots advance in high-risk environments, such as the new Mars Perseverance Rover [1]. This new rover differs from its predecessor, Mars Curiosity Rover, in that it has the independence to cover ground and make decisions without being directly controlled by its human operators on Earth. It also has planning features that allow it to shift its daily activities around to be more efficient with openings in its daily schedule. Another example is urban search and rescue (USAR) robots [2]. These robots are deployed in damaged buildings that are too dangerous for human rescuers to enter. The robots are responsible for mapping the environment, locating victims, and navigating the environment. The Mars Perseverance and Curiosity Rovers and USAR robots are just a few examples out of many of advanced autonomous technology that provides opportunities for studying how these advanced technologies might assist, collaborate, and team with humans.

A *team* can be defined as a sociotechnical system that contains two or more heterogeneous and interdependent members (either human or nonhuman, e.g., machines and canines) who interact with each other to complete a common goal or task [3]. Therefore, there are also other mixed-teams,

such as human-machine teaming (HMT) [4]. HMT is a sociotechnical system in which at least one machine act as heterogeneous and interdependent team members and interact with human team members to accomplish a common goal or task [5]. In order to design better HMTs, it is important to understand their team processes, such as team communication and trust. Consequently, the main focus of the current study is to empirically examine the relationship between communication and trust in an HMT within a simulated remotely piloted aircraft system (RPAS) context.

II. HUMAN-MACHINE TEAMING

HMT is a broad term that we apply to a specific team definition that considers task role heterogeneity, interdependence, and a common team goal or task. This section clarifies the machine concept and how it can be a teammate within this definition. To understand how people might team with autonomous technology and how trust might be an important construct to consider when studying HMT, it is important to understand what autonomous technology is (traditionally referred to as ‘machine’) and how a machine can be classified as automation or autonomy [6]–[8]. A machine is a “device, having a unique purpose, that augments or replaces human or animal effort for the accomplishment of physical tasks” [9]. Sheridan (2002) underlines a three-part definition of what *automation* is: “(1) the mechanization and integration of the sensing of environmental variables (by artificial sensors); (2) data processing and decision making (by computers); and (3) mechanical action (by motors or devices that apply forces on the environment)” [11; p.9]. On the other hand, *autonomy* can be thought of as a machine that can carry out tasks independently or in conjunction with human input and oversight beyond that of what is traditionally considered “automation” [6].

Beyond definitions, though, the difference between automation and autonomy can be thought of as a spectrum. In this spectrum, *automation* is on the low end where the technology is not autonomous and requires human oversight and intervention, whereas *autonomy* is on the high end and is technology that is independent of human input and oversight [11]–[14]. Accordingly, the overlap with all of these taxonomies seems to be that: (1) on the low end, the human does everything; (2) in the middle, the autonomous technology carries out a task or informs the human of certain variables to help the human make a decision; and (3) at the higher levels, the autonomous

technology has complete autonomy and does not require the permission of humans to carry out tasks, but can work with the human in a team like setting.

Autonomy is becoming more sophisticated and better at complementing humans in team-like settings by doing things that the humans cannot do or prefer not to do. And as autonomy continues this trend, it might mean that autonomy might work with humans as a teammate and an independent entity fulfilling a role not completed by any other teammate(s) on the team [6]. This is important, especially in the context of trust, because in a team, team members have *heterogeneous* roles, meaning that each team member has a specific task and is responsible for executing that task within the team [15]. If a teammate cannot be trusted to fulfill their role, the team cannot achieve its overall goals effectively. Furthermore, an autonomous teammate is *interdependent* with other teammates, which helps the team achieve their overall goals [6], [16], [17]. Research has indicated that interdependence with an autonomous agent helps human teammates perceive the autonomous agent as more cooperative, friendlier, and as if the technology provided the human teammate with high-quality information [18]. Although these outcomes might seem beneficial in the sense that they will make the human trust the autonomy more, increased trust does not necessarily equate with appropriately calibrated trust.

III. TRUST IN HUMAN-MACHINE TEAMING

Trust is another critical dimension in effective HMT. One commonly accepted definition of trust is “the willingness of a party to be vulnerable to the actions of another party based on the expectation that the other will perform a particular action important to the trustor, irrespective of the ability to monitor or control that other party” [19, p. 712]. This definition alone, however, does not capture the essence of what trust is. To fully understand trust as a cognitive process, it is important to understand the structural elements of trust.

One of the structural elements of trust is interpersonal trust, which is “the generalized expectancy held by an individual that the word, promise, oral or written statement of another individual or group can be relied on” [20, p. 651]. Risk is the main factor in understanding interpersonal trust. Rousseau and colleagues (1998) argued that uncertainty is the source of risk, and only when there is risk can one trust. When someone chooses to trust, they engage in risk-taking behavior. This aligns with another study that indicates that trust is not actually taking a risk, but rather it is a willingness to take a risk [19]. That is, the willingness to take a risk (trust) precedes the actual taking of risk (trusting behavior); trust (or trusting) is not a behavior but rather leads to behavior. Therefore, in order for one party to truly trust another party, the other party must be seen as trustworthy.

Trustworthiness and interpersonal trust are integrated; interpersonal trust is the belief that someone will do what they say they will do [20]. Interpersonal trust, therefore, not only serves as a foundation for being perceived as initially trustworthy, but also for being perceived as trustworthy in the future. Indeed, trustworthiness is time-based. For instance, one’s reputation is based on actions and behaviors that one has done in the past. If those actions and behaviors lead to one having a positive reputation, the individual will be perceived as trustworthy. This perceived trustworthiness based on a person’s

reputation will lead to others initially having high interpersonal trust towards that person even though they may not have previously interacted with that person [21]. Although risk and trustworthiness are indeed structural elements that can explain the nature of trust, they are limited in their capacity to define it for trust in a machine teammate in HMT.

Trust in a machine teammate is a complex idea that has many dimensions worth considering, such as trust calibration between two team members and trust in other team members. Autonomy can be considered as a team member because of the conceptual similarities of “teams” and “autonomy” as seen in the aforementioned definitions provided in this paper (e.g., interdependence, task completion, interaction/communication, etc.). Chen and Barnes (2014) describe multiple factors that influence trust in autonomous teammates, such as having a shared cognitive architecture; so that the beliefs, desires, and intentions of the humans and the autonomy are compatible [22]. Demir et al. (2021) discuss this multidimensional perspective in the taskwork and teamwork of HMT [23]. The results of their study were mixed in that the researchers found stable team interactions to be negatively associated with trust development, but beyond an inflection point, they were positively associated with trust development. Results also showed that recovery from autonomy failures was related to a moderate amount of trust, but too little or too much trust led to poorer recovery from autonomy failures. These results indicate that trusting an autonomous teammate is somewhat linked to team interactions, and it is an evolving process. In this study, we also consider trust as an evolving process that is related to team communication.

IV. COMMUNICATION IN HUMAN-MACHINE TEAMING

Behavioral and physiological measures play an important role in assessing trust from a dynamic perspective. Several studies [23]–[25] indicate that team interaction based on communication and coordination can give a bigger picture of trust in the light of the theory of interactive team cognition (ITC; [26]). Interactions are vital to team success and performance. Because of this, in order to reflect interactions within a team, Cooke et al., 2009 developed a communication analysis that looks at who is talking to whom, what is being said (behavior/ content), and communication flow [15].

In this study, we focus on specific types of communication, i.e., anticipatory pushing of information. We define anticipatory pushing as pushing information from Teammate-A to Teammate-B based on Teammate-A’s anticipation that Teammate-B requires the information. It is hypothesized that the anticipation of information for a teammate stems from a good shared mental model among teammates that leads to implicit coordination (i.e., coordination that is not preplanned or explicitly communicated but rather arises out of the shared understanding of a given situation) [27]. Anticipatory pushing of information is an example of implicit coordination because one team member is anticipating the informational needs of another teammate helping another teammate, perhaps by adjusting to a change by notifying them of the change in the situation and possibly recommending an action that addresses the change in the situation [6].

With that in mind, the following novel definition of trust in a teammate based on anticipation is proposed for the current study. *Trust in a teammate is the expectation that teammates will share anticipated and needed information with each other such that the sharing of information facilitates goal accomplishment, safety, and task continuation and completion.* Furthermore, this definition of trust is context-free and can be applied to any team. Indeed, regardless of context, all teams need to share information to accomplish their overarching goal(s), which is why the team was formed in the first place. Teammates need to keep each other safe, not only physically but also mentally and emotionally. Sharing information based on anticipated needs certainly can help with this goal of teammate preservation. And finally, information sharing is necessary for task continuation, as task continuation and completion help the team accomplish its goal(s) in a timely manner. Because the current study is focused on information sharing as to how trust is formed, maintained, and calibrated, a corresponding measure of trust has to be implemented.

V. CURRENT STUDY

A. Synthetic Team Task and Roles

The experiment took place in the RPAS-Synthetic Task Environment (RPAS-STE) testbed, which mimics the individual and team cognitive tasks in an RPAS ground station. The RPAS ground station comprises three heterogeneous and interdependent task roles (Fig. 1 for the role definitions; [28]). The task was to take good photos of targets while navigating the remotely piloted aircraft (RPA) along a safe route.

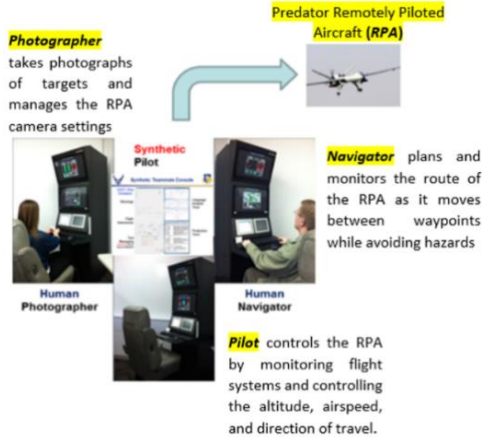


Fig. 1. The Ground Station Consoles in RPAS-STE [30]. Each of the three-team members communicated via a touch-screen text-chat interface.

In this study, the pilot role was an “AI” teammate that was simulated by a trained confederate utilizing the “Wizard of Oz” methodology (WoZ) [29]. The confederate followed a script that indicated when and what to communicate throughout the task and described behaviors for controlling the flight of the RPA. This allowed us to consistently and reliably introduce the same autonomy failures across teams. For this experiment, the capabilities of the AI were limited in verbal comprehension, production, and piloting behaviors. Misspellings and unclear

communications were ignored, and the pilot’s behaviors were generally limited to the script.

B. Design and Research Question

This study consists of between-and within-subjects design manipulations. However, we only consider the within-subjects design manipulations because of the main focus of the current study and page limit. There are three between-subjects design effects based on the pre-mission training, including coordination training, calibration training, and control training—all defined by manipulating pre-mission training [31]. The within-subjects design includes routine conditions (i.e., missions with no technology failures) and degraded conditions (i.e., missions with technology failures; see Table I). Because the focus of this study was to identify the relationship between anticipatory pushing of information and trust under routine and degraded conditions in HMT, we combined all six types of technology failures together into one condition (see Table I [31]). This allowed us to analyze the data more simply and match the dimension of the trust measure, which was obtained via questionnaire once following routine conditions and again following degraded conditions (as within-subjects design).

TABLE I. FAILURES WITHIN THE DEGRADED CONDITION [31].

Type	Description
Automation	Prevented display of flight information such as airspeed, altitude, or heading to the photographer or pilot.
Automation	A one-way communication cut between photographer and pilot.
Automation	A gradual power-down and subsequent power-up of all six workstation screens, affecting all experimental positions.
Autonomy	Simulated a malfunction in the AI teammate’s capacity for properly responding to messages from teammates.
Autonomy	Simulated a hijacking of the RPAS by moving it to an enemy waypoint while the AI agent provided deceptive responses.
Hybrid	A combination of both automation and autonomy failures into a single failure.

The current study focuses on trust in a teammate, regardless of whether a teammate is a human or machine, and how to measure trust. The goal is to assess whether a proposed trust metric is a valid metric for measuring trust in a teammate, using interactions as measurement units. The specific interactions being used as the units of measure are interactions where information pushing based on anticipated needs occurs. It is assumed that as a teammate anticipates another teammate’s needs and pushes information to that teammate based on anticipated needs, there will be increased trust towards the teammate who pushed the information. This will lead to higher levels of interpersonal trust towards that teammate.

VI. METHOD

A. Participants

Sixty-four participants were recruited from Arizona State University and surrounding areas. Participants ranged in age from 18 to 33 ($M_{\text{age}} = 22.53$, $SD_{\text{age}} = 3.55$). A total of 60 participants, divided into 30 teams, completed the study. Participants were randomly assigned to the roles of *navigator*

or *photographer*. An experimenter filled the pilot role. Each team completed approximately a seven-hour session. Participants were compensated \$10 per hour for participation.

B. Procedure

After providing their informed consent, participants were randomly assigned to their respective workstations separated by a partition. The confederate pilot was located in a different room as a part of the WoZ methodology. The participants then completed training according to their assigned roles. Training consisted of an interactive slideshow that described their roles and tasks and a 40-minute hands-on training mission that familiarized the participants with the interfaces, roles, and communication in the RPAS task environment. A trained experimenter guided them using a script to ensure the participants understood the task. Participants were also aware of a checklist relating to their roles and guidance for communicating with the AI teammate. The first mission was a baseline mission with no failures. Each mission consisted of 11–20 targets. Short breaks (15min) were given to participants between each mission. Following the first and fifth mission, questionnaires assessing trust and workload were administered. After the fifth mission, trust and demographics questionnaire was also completed, and participants were debriefed.

TABLE II. EXPERIMENTAL SESSION AND FAILURES

<i>Session Order</i>	<i>Failure I</i>	<i>Failure II</i>
Mission 1 (Routine)	No Failure	No Failure
Pre-Questionnaires (Trust Questionnaire)		
Mission 2 (Degraded)	Automation	Autonomy
Mission 3 (Degraded)	Autonomy	Automation
Mission 4 (Degraded)	Hybrid	Automation
Mission 5 (Degraded)	Automation	Autonomy
Pre-Questionnaires (Trust Questionnaire)		

C. Measures

In this study, we collected several measures, including individual and team performance scores, team situation awareness, process measures (team communication behaviors and flow, coordination, process ratings, sensor-based metrics (electrocardiogram and facial expressions), NASA Task Load Index [32], trust [33], and demographic questions. To address the research questions, we only considered the measures of (1) *trust* and (2) *anticipatory pushing*. *Trust* was measured based on the questionnaire used by [30], a modified version of the [19] questionnaire. The questionnaire used by [30] consisted of 18 items (9 items per teammate) and used a scoring scale of one to five. To obtain the means of the participants' reported trust score, 4 of the 18 items were reverse-scored to align with the scale for the remaining 14 questions.

Anticipatory pushing refers to text-chat data that shows the pushing of information from one teammate to another without being explicitly asked. It is important to note here that because the AI teammate essentially followed a script and only answered questions pertaining to its role and responsibilities, it

did not do any anticipatory pushing of information to either human teammate in either of the two conditions. Two experimenters recorded communication behaviors. Inter-rater reliability was assessed for agreement between experimenters when recording anticipatory pushing of information. Fleiss' Kappa showed a good agreement between the experimenters' judgments, $\kappa=0.869$ (95% CI, 0.842 to 0.895), $p < 0.0001$.

VII. DATA ANALYTICS AND RESULTS

A Multivariate Analysis of Variance (MANOVA) was applied to address the following questions: (1) How does anticipatory pushing of information and trust change from human to human and human to autonomous team member across the two sessions (i.e., routine and degraded)?, and (2) How is anticipatory pushing of information associated with trust in a teammate across the two sessions? The test statistics show that all the effects were statistically significant (Table III). Mauchly's test indicated that the assumption of sphericity were not satisfied for pair [$\chi^2(5)=41.3$, $p < 0.001$, $\epsilon = 0.501$], factor by pair [$\chi^2(5)=65.4$, $p < 0.001$, $\epsilon = 0.435$], pair by session [$\chi^2(5)=48.02$, $p < 0.001$, $\epsilon = 0.506$], and factor by pair by session [$\chi^2(5)=52.8$, $p < 0.001$, $\epsilon = 0.457$]. Therefore, degrees of freedom were corrected using the Greenhouse-Geisser correction for within-subjects effects. Accordingly, all the within-subjects effects were statistically significant (Table IV).

TABLE III. MULTIVARIATE TEST STATISTICS

<i>Effect</i>	<i>Wilk's A</i>	<i>df_{Hypothesis}</i>	<i>df_{Error}</i>	<i>p-value</i>	<i>η_p^2</i>
Factor	0.084	1	27	0.000	0.916
Pair	0.480	3	25	0.000	0.520
Session	0.670	1	27	0.001	0.330
Factor by Pair	0.447	3	25	0.000	0.553
Factor by Session	0.567	1	27	0.000	0.433
Pair by Session	0.562	3	25	0.002	0.438
Factor by Pair by Session	0.469	3	25	0.000	0.531

Based on the significant interaction effect of factor by pair and by session, we evaluated Least Significant Difference (LSD) pairwise comparisons. First, we evaluated interaction effects for anticipatory pushing of information across the pairs and sessions (see Fig. 2). Across all the pairs (except from the photographer to the pilot, $p = 0.060$), anticipatory pushing of information was significantly higher in the degraded conditions than the routine conditions ($p < 0.05$). Also, for all of the pairs, the highest pushing of information occurred from the navigator to the pilot ($p < 0.001$). Due to the WoZ manipulation, there was no anticipatory pushing of information from the AI teammate to a human team member. The navigator did more anticipatory pushing of information to the photographer and pilot ($p < 0.05$) than the other pairs. The overall findings indicate that human team members pushed more information during the degraded conditions and demonstrated more control over the AI teammate because of its abnormal behaviors (i.e., autonomy failures), perhaps viewing it more as automation and not an autonomous team member they were interdependent with. The human team members in the routine conditions

anticipated each other's needs significantly less than the needs of the AI team member ($p < 0.05$). A possible reason might be that the pilot role of the AI was central in interaction because it needed information from both human teammates in order to complete each task.

TABLE IV. TEST OF WITHIN-SUBJECTS EFFECTS

Effect	$df_{Hypothesis}$	df_{Error}	F -Test	p -value	η_p^2
Factor	1	27	294.65	0.000	0.916
Pair	1.51	81	14.11	0.000	0.343
Session	1	27	13.31	0.001	0.330
Factor by Pair	1.31	81	21.61	0.000	0.445
Factor by Session	1	27	20.62	0.000	0.433
Pair by Session	1.52	81	5.72	0.011	0.175
Factor by Pair by Session	1.37	81	6.49	0.009	0.194

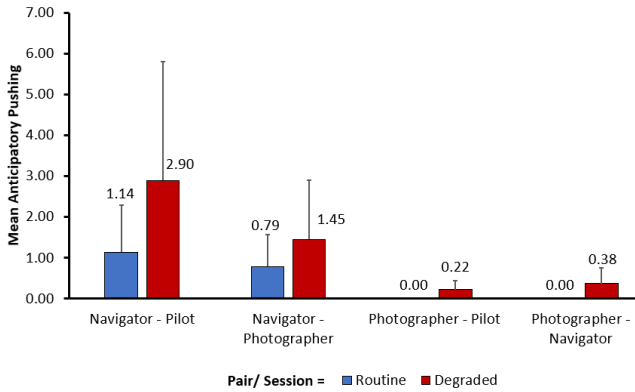


Fig. 2. Anticipatory pushing of information across the pairs and sessions.

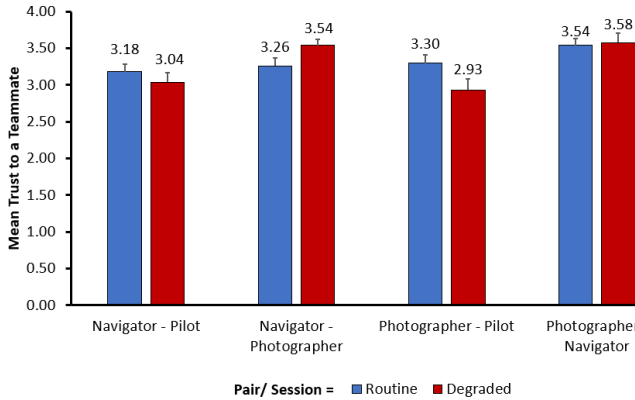


Fig. 3. Trust in a teammate across the pairs and sessions.

However, the findings for trust contrasted with the anticipatory pushing of information (Fig. 3). In degraded conditions, trust in the AI team member was associated with significantly lower trust in a human team member by the same individual ($p < 0.05$). This was especially seen under degraded conditions in comparison to routine conditions. Another finding showed that trust from the navigator to the photographer was significantly higher in routine conditions as compared to degraded conditions ($p < 0.05$). This finding makes sense when it is compared with the pushing of information. Anticipatory

pushing of information increased from the navigator to the photographer and vice versa over time ($p < 0.05$). Overall, these findings indicate that anticipatory pushing of information was associated with increased trust between the human teammates but not trust from humans to autonomy.

VIII. DISCUSSION AND CONCLUSION

Effective human-machine teaming requires that teammates anticipate one another's needs, share information proactively, and have appropriate trust. This research explored the relationship between anticipatory pushing of information and trust in a teammate in a remotely piloted aircraft system synthetic task environment. Teams of two participants and one AI teammate(emulated by a confederate) executed a series of missions under routine and degraded conditions, and measures of anticipatory pushing and trust were analyzed.

Our first research question was concerned with how trust and anticipatory pushing compared between human-human dyads and human-autonomy dyads in routine and degraded conditions. Our finding indicated that the human participants exhibited more anticipatory pushing behaviors when conditions were degraded compared to when they were not. These pushes were directed to both their human and machine counterparts. Anticipatory pushing of information has been associated with the development of more implicit coordination strategies as teams become more familiar [27], and may have helped teams compensate when task load was increased due to degraded conditions and coordination costs were higher. The highest number of pushes was observed from the navigator role to the pilot. This demonstrates the importance of how the availability and interdependence of information provided to operators may play a role in both their individual taskwork and also how they coordinate with others in the HMT. In this case, the pilot may have required more direction under degraded conditions, but the navigator was often the only one who could provide it. We also found that trust in the AI teammate appeared to diminish when conditions were degraded, whereas trust in human teammates tended to increase. The reduction in trust in the AI teammate may have been due to the failures of the autonomous agent that resulted in undesirable behaviors (autonomy failures). Research suggests that trust violations are treated differently when a team member is an artificial agent, and that trust repair strategies can improve performance [7]. In contrast, trust between the human teammates may have evolved over time due to shared experiences and positive interactions over the course of the study. This suggests that degraded conditions may impact trust dynamics in HMTs differently, depending on the makeup of the team.

Our second research question was concerned with whether anticipatory pushing of information was associated with trust. Our findings suggest that anticipatory pushing of information was associated with increased trust between the human teammates, but not trust from humans to autonomy. However, the direction of causality is not clear. Trust may increase anticipatory pushing behaviors. Trust is probably necessary to enable effective coordination and communication. Or vice

versa, anticipatory pushing of information may increase trust. Therefore, it is possible that anticipatory pushing of information could be used as a metric for measuring trust. Future research looking at how varying levels of anticipatory pushing correlate with self-reported trust scores could be conducted to validate this claim. Additionally, it might be possible to determine how anticipatory pushing of information as a metric and trust as a team concept relate to team performance. Previous research has shown that communication and trust are essential components of good team performance.

There were several limitations in this study. One is the linear and nonlinear relationship between anticipatory pushing and trust. It is possible that there is a positive correlation between trust and anticipatory pushing of information up to a point, but then an inflection point is reached where there is a negative correlation. Furthermore, we did not consider anticipatory pushing from the autonomous agent to the human participants. Researchers in the future might be interested in employing a WoZ methodology to determine this. Finally, these results also need to be evaluated in between-subjects effects (i.e., training).

ACKNOWLEDGMENT

This research is supported by ONR Award N000141712382 (Program Managers: Marc Steinberg, Micah Clark).

REFERENCES

- [1] mars.nasa.gov, "Mars 2020 Perseverance Rover." <https://mars.nasa.gov/mars2020/> (accessed May 07, 2021).
- [2] R. R. Murphy *et al.*, "Search and Rescue Robotics," in *Springer Handbook of Robotics*, B. Siciliano and O. Khatib, Eds. Berlin: Springer, 2008, pp. 1151–1173. doi: 10.1007/978-3-540-30301-5_51.
- [3] J. A. Cannon-Bowers, E. Salas, and S. A. Converse, "Cognitive psychology and team training: Training shared mental models and complex systems," *Human factors society*, vol. 33, no. 12, pp. 1–4, 1990.
- [4] M. Demir and N. J. Cooke, "Human Teaming Changes Driven by Expectations of a Synthetic Teammate," *Proceedings of the Human Factors and Ergonomics Society Annual Meeting*, vol. 58, no. 1, pp. 16–20, Sep. 2014, doi: 10.1177/1541931214581004.
- [5] M. Demir, N. J. McNeese, and N. J. Cooke, "Team situation awareness within the context of human-autonomy teaming," *Cognitive Systems Research*, vol. 46, pp. 3–12, 2017, doi: 10.1016/j.cogsys.2016.11.003.
- [6] N. J. McNeese, M. Demir, N. J. Cooke, and C. Myers, "Teaming With a Synthetic Teammate: Insights into Human-Autonomy Teaming," *Hum Factors*, vol. 60, no. 2, 2018, doi: 10.1177/0018720817743223.
- [7] E. J. de Visser, R. Pak, and T. H. Shaw, "From 'automation' to 'autonomy': The importance of trust repair in human-machine interaction," *Ergonomics*, 2018, doi: 10.1080/00140139.2018.1457725.
- [8] D. B. Kaber, "A conceptual framework of autonomous and automated agents," *Theoretical Issues in Ergonomics Science*, vol. 19, no. 4, pp. 406–430, Jul. 2018, doi: 10.1080/1463922X.2017.1363314.
- [9] "Machine," *Encyclopedia Britannica*. <https://www.britannica.com/technology/machine> (accessed May 07, 2021).
- [10] T. B. Sheridan, *Humans and automation: system design and research issues*. Hoboken, N.J.: Wiley [u.a.], 2002.
- [11] M. R. Endsley and D. B. Kaber, "Level of automation effects on performance, situation awareness and workload in a dynamic control task," *Ergonomics*, Mar. 1999, doi: 10.1080/001401399185595.
- [12] M. R. Endsley and E. O. Kiris, "The Out-of-the-Loop Performance Problem and Level of Control in Automation," *Human Factors: The Journal of the Human Factors and Ergonomics Society*, vol. 37, no. 2, pp. 381–394, Jun. 1995, doi: 10.1518/001872095779064555.
- [13] R. Parasuraman, T. B. Sheridan, and C. D. Wickens, "A model for types and levels of human interaction with automation," *IEEE Transactions on Systems, Man and Cybernetics, Part A: Systems and Humans*, vol. 30, no. 3, pp. 286–297, May 2000, doi: 10.1109/3468.844354.
- [14] T. B. Sheridan and W. L. Verplank, "Human and Computer Control of Undersea Teleoperators," Massachusetts Inst of Tech Cambridge Man-Machine Systems Lab, Jul. 1978. Accessed: Apr. 30, 2020. [Online]. Available: <https://apps.dtic.mil/docs/citations/ADA057655>
- [15] N. J. Cooke and J. C. Gorman, "Interaction-Based Measures of Cognitive Systems," *Journal of Cognitive Engineering and Decision Making*, vol. 3, no. 1, pp. 27–46, Mar. 2009, doi: 10.1518/155534309X433302.
- [16] J. B. Lyons, K. T. Wynne, S. Mahoney, and M. A. Roebke, "Chapter 6 - Trust and Human-Machine Teaming: A Qualitative Study," in *Artificial Intelligence for the Internet of Everything*, W. Lawless, R. Mittu, D. Sofge, I. S. Moskowitz, and S. Russell, Eds. Academic Press, 2019, pp. 101–116. doi: 10.1016/B978-0-12-817636-8.00006-5.
- [17] K. T. Wynne and J. B. Lyons, "An integrative model of autonomous agent teammate-likeness," *Theoretical Issues in Ergonomics Science*, vol. 19, no. 3, pp. 353–374, May 2018, doi: 10.1080/1463922X.2016.1260181.
- [18] C. Nass, B. J. Fogg, and Y. Moon, "Can computers be teammates?," *International Journal of Human-Computer Studies*, vol. 45, no. 6, pp. 669–678, Dec. 1996, doi: 10.1006/ijhc.1996.0073.
- [19] R. C. Mayer, J. H. Davis, and F. D. Schoorman, "An Integrative Model of Organizational Trust," *The Academy of Management Review*, vol. 20, no. 3, pp. 709–734, 1995, doi: 10.2307/258792.
- [20] J. B. Rotter, "A new scale for the measurement of interpersonal trust1," *J Personality*, vol. 35, no. 4, pp. 651–665, Dec. 1967, doi: 10.1111/j.1467-6494.1967.tb01454.x.
- [21] D. H. McKnight, L. L. Cummings, and N. L. Chervany, "Initial Trust Formation in New Organizational Relationships," *The Academy of Management Review*, vol. 23, no. 3, 1998, doi: 10.2307/259290.
- [22] J. Y. C. Chen and M. J. Barnes, "Human-Agent Teaming for Multirobot Control: A Review of Human Factors Issues," *IEEE Transactions on Human-Machine Systems*, 2014, doi: 10.1109/THMS.2013.2293535.
- [23] M. Demir, N. J. McNeese, J. C. Gorman, N. J. Cooke, C. W. Myers, and D. A. Grimm, "Exploration of Team Trust and Interaction Dynamics in Human-Autonomy Teaming," *IEEE on Human-Machine Sys*, 2021.
- [24] N. Tenhundfeld, M. Demir, and E. J. de Visser, "An Argument for Trust Assessment in Human-Machine Interaction: Overview and Call for Integration." OSF Preprints, Jan. 11, 2021, doi: 10.31219/osf.io/j47df.
- [25] L. Huang *et al.*, "Chapter 13 - Distributed dynamic team trust in human, artificial intelligence, and robot teaming," in *Trust in Human-Robot Interaction*, C. S. Nam and J. B. Lyons, Eds. Academic Press, 2021, pp. 301–319. doi: 10.1016/B978-0-12-819472-0.00013-7.
- [26] N. J. Cooke, J. C. Gorman, C. W. Myers, and J. L. Duran, "Interactive Team Cognition," *Cognitive Science*, vol. 37, no. 2, pp. 255–285, Mar. 2013, doi: 10.1111/cogs.12009.
- [27] E. E. Entin and D. Serfaty, "Adaptive Team Coordination," *Human Factors*, vol. 41, no. 2, 1999, doi: 10.1518/001872099779591196.
- [28] N. J. Cooke and S. M. Shope, "Designing a Synthetic Task Environment," in *Scaled Worlds: Development, Validation, and Application*, L. R. E. Schiflett, E. Salas, and M. D. Coovert, Eds. Surrey, England: Ashgate Publishing, 2004, pp. 263–278.
- [29] L. D. Riek, "Wizard of Oz Studies in HRI: A Systematic Review and New Reporting Guidelines," *J. Hum.-Robot Interact.*, vol. 1, no. 1, pp. 119–136, Jul. 2012, doi: 10.5898/JHRI.1.1.Riek.
- [30] C. J. Johnson *et al.*, "Training and Verbal Communications in Human-Autonomy Teaming Under Degraded Conditions," in *2020 IEEE COGSIMA*, 2020, doi: 10.1109/CogSIMA49017.2020.9216061.
- [31] M. Demir, N. J. McNeese, C. Johnson, J. C. Gorman, D. Grimm, and N. J. Cooke, "Effective Team Interaction for Adaptive Training and Situation Awareness in Human-Autonomy Teaming," in *2019 IEEE COGSIMA*, 2019, doi: 10.1109/COGSIMA.2019.8724202.
- [32] S. G. Hart and L. E. Staveland, "Development of NASA-TLX (Task Load Index): Results of empirical and theoretical research," in *Human Mental Workload*, P. A. Hancock and N. Mashkati, Eds, 1988, pp. 139–183.
- [33] R. C. Mayer and M. B. Gavin, "Trust in Management and Performance: Who Minds the Shop while the Employees Watch the Boss?," *The Academy of Management Journal*, vol. 48, no. 5, pp. 874–888, 2005, doi: 10.2307/20159703.