

**DEVELOPING OBJECTIVE COMMUNICATION-BASED
MEASURES OF TRUST FOR HUMAN-AUTONOMY TEAMS**

A Thesis
Presented to
The Academic Faculty

by

Matthew J. Scalia

In Partial Fulfillment
of the Requirements for the Degree
Master's of Science in the
School of Psychology

Georgia Institute of Technology
August 2022

COPYRIGHT © 2022 BY MATTHEW J. SCALIA

**DEVELOPING OBJECTIVE COMMUNICATION-BASED
MEASURES OF TRUST FOR HUMAN-AUTONOMY TEAMS**

Approved by:

Dr. Jamie Gorman, Advisor
School of Psychology
Georgia Institute of Technology

Dr. Bruce Walker
School of Psychology
Georgia Institute of Technology

Dr. Christopher Wiese
School of Psychology
Georgia Institute of Technology

Date Approved: July 14, 2022

To my family, Judy, Frank, and Andrew Scalia. Thank you for your unwavering love and support.

ACKNOWLEDGEMENTS

Thank you to my advisor, Dr. Jamie Gorman, and my committee, Dr. Bruce Walker and Dr. Christopher Wiese, for their guidance throughout this project. Thank you to my lab mates David Grimm and Shiwen Zhou who helped conduct the experiment for this project. Further, I would like to thank all my lab mates—Terri Dunbar, David Grimm, Julie Harrison, and Shiwen Zhou—for providing feedback on my ideas, presentations, and for their encouragement. Thanks to research assistants Anna Crofton, Fiorella Gambetta, Yiwen Zhao, and Ansley Lee who provided their time to run experimental sessions and code data. Finally, thanks to my fellow graduate students for their support in all my endeavors.

TABLE OF CONTENTS

ACKNOWLEDGEMENTS	iv
LIST OF TABLES	vii
LIST OF FIGURES	xi
LIST OF SYMBOLS AND ABBREVIATIONS	xii
SUMMARY	xiii
CHAPTER 1. Introduction	1
1.1 The Current Study	7
1.2 Hypotheses	8
1.2.1 Hypothesis 1	9
1.2.2 Hypothesis 2	9
CHAPTER 2. Method	10
2.1 Participants	10
2.2 Materials	10
2.3 Procedure	12
2.4 Measures	13
2.4.1 ICM	13
2.4.2 Team Trust	14
2.4.3 Team Performance	16
2.5 Design	16
2.5.1 The Third Variable Problem	17
CHAPTER 3. Results	20
3.1 Exploratory Factor Analyses	20
3.1.1 Exploratory Factor Analysis in All-Human Condition	20
3.1.2 Exploratory Factor Analysis in Human-Autonomy Condition	22
3.2 Third Variable Models	24
3.2.1 ICM, Team Trust, and Team Performance in All-Human Teams	24
3.2.2 ICM, Team Trust, and Team Performance in HATs	32
CHAPTER 4. Discussion	41
4.1 Result Summary of Third Variable Analyses for All-Human and HATs	41
4.2 Result Summary of Third Variable Analyses for All-Human and HATs Using Session 1 and Session 2 Team Trust Scores	42
4.3 Summary of Results for the ICM Component Regressions	44
4.4 Aspects of Team Trust in All-Human Teams and HATs	46
4.5 Limitations and Future Directions	49
4.6 Conclusion	51
APPENDIX A. Modified Trust Questionnaire	53

APPENDIX B. Checklist for Trust between People and Automation Scale	56
REFERENCES	58

LIST OF TABLES

Table 1	– Experimental Session	12
Table 2	– Team Verbal Behaviors	14
Table 3	– Team Verbal Behavior Inter-Rater Reliability Cohen’s K	14
Table 4	– Factor Loadings on the First Three Factors in the All-Human Condition	21
Table 5	– Factor Loadings on the First Three Factors in the Human-Autonomy Condition	23
Table 6	– Summary of Regression Analysis for Model 1 in the All-Human Condition	24
Table 7	– Summary of Regression Analysis for Model 2 in the All-Human Condition	24
Table 8	– Summary of Regression Analysis for Model 3 in the All-Human Condition	25
Table 9	– Summary of Regression Analysis for Model 1 using Team Trust Scores from Session 1 in the All-Human Condition	25
Table 10	– Summary of Regression Analysis for Model 2 using Team Trust Scores from Session 1 in the All-Human Condition	25
Table 11	– Summary of Regression Analysis for Model 3 using Team Trust Scores from Session 1 in the All-Human Condition	26
Table 12	– Summary of Regression Analysis for Model 1 using Team Trust Scores from Session 2 in the All-Human Condition	26
Table 13	– Summary of Regression Analysis for Model 2 using Team Trust Scores from Session 2 in the All-Human Condition	27
Table 14	– Summary of Regression Analysis for Model 3 using Team Trust Scores from Session 2 in the All-Human Condition	27
Table 15	– Summary of Regression Analysis for Average Amount of Messages Predicting Team Trust in the All-Human Condition	28

Table 16	– Summary of Regression Analysis for Average Amount of Messages Predicting Team Performance in the All-Human Condition	28
Table 17	– Summary of Regression Analysis for Team Message Frequency Predicting Team Trust in the All-Human Condition	28
Table 18	– Summary of Regression Analysis for Team Message Frequency Predicting Team Performance in the All-Human Condition	29
Table 19	– Summary of Regression Analysis for Team Affect Predicting Team Trust in the All-Human Condition	29
Table 20	– Summary of Regression Analysis for Team Affect Predicting Team Performance in the All-Human Condition	30
Table 21	– Summary of Regression Analysis for Team Pushing Verbal Behavior Predicting Team Trust in the All-Human Condition	30
Table 22	– Summary of Regression Analysis for Team Pushing Verbal Behavior Predicting Team Performance in the All-Human Condition	31
Table 23	– Summary of Regression Analysis for Team Pulling Verbal Behavior Predicting Team Trust in the All-Human Condition	31
Table 24	– Summary of Regression Analysis for Team Pulling Verbal Behavior Predicting Team Performance in the All-Human Condition	32
Table 25	– Summary of Regression Analysis for Model 1 in the Human-Autonomy Condition	32
Table 26	– Summary of Regression Analysis for Model 1 in the Human-Autonomy Condition	32
Table 27	– Summary of Regression Analysis for Model 2 in the Human-Autonomy Condition	33
Table 28	– Summary of Regression Analysis for Model 1 using Team Trust Scores from Session 1 in the Human-Autonomy Condition	33
Table 29	– Summary of Regression Analysis for Model 2 using Team Trust Scores from Session 1 in the Human-Autonomy Condition	33
Table 30	– Summary of Regression Analysis for Model 3 using Team Trust Scores from Session 1 in the Human-Autonomy Condition	34
Table 31	– Summary of Regression Analysis for Model 1 using Team Trust Scores from Session 2 in the Human-Autonomy Condition	34

Table 32	– Summary of Regression Analysis for Model 2 using Team Trust Scores from Session 2 in the Human-Autonomy Condition	35
Table 33	– Summary of Regression Analysis for Model 3 using Team Trust Scores from Session 2 in the Human-Autonomy Condition	35
Table 34	– Summary of Regression Analysis for Average Amount of Messages Predicting Team Trust in the Human-Autonomy Condition	35
Table 35	– Summary of Regression Analysis for Average Amount of Messages Predicting Team Performance in the Human-Autonomy Condition	36
Table 36	– Summary of Regression Analysis for Team Message Frequency Predicting Team Trust in the Human-Autonomy Condition	36
Table 37	– Summary of Regression Analysis for Team Message Frequency Predicting Team Performance in the Human-Autonomy Condition	37
Table 38	– Summary of Regression Analysis for Team Affect Predicting Team Trust in the Human-Autonomy Condition	37
Table 39	– Summary of Regression Analysis for Team Affect Predicting Team Performance in the Human-Autonomy Condition	38
Table 40	– Summary of Regression Analysis for Team Pushing Verbal Behavior Predicting Team Trust in the Human-Autonomy Condition	38
Table 41	– Summary of Regression Analysis for Team Pushing Verbal Behavior Predicting Team Performance in the Human-Autonomy Condition	38
Table 42	– Summary of Regression Analysis for Team Pulling Verbal Behavior Predicting Team Trust in the Human-Autonomy Condition	39
Table 43	– Summary of Regression Analysis for Team Pulling Verbal Behavior Predicting Team Performance in the Human-Autonomy Condition	39
Table 44	– Descriptive Statistics of All-Human Teams' Trust Change Scores at the Individual Level	42
Table 45	– Descriptive Statistics of HATs' Trust Change Scores at the Individual Level	42
Table 46	– Modified Trust Questionnaire	53

LIST OF FIGURES

Figure 1	– Third variable model	17
Figure 2	– Model 1 for both conditions.	18
Figure 3	– Model 2 for both conditions.	18
Figure 4	– Model 3 for both conditions.	19
Figure 5	– All-human scree plot	20
Figure 6	– HAT scree plot	22

LIST OF SYMBOLS AND ABBREVIATIONS

CERTT-	Cognitive Engineering Research on Team Tasks Remote Piloted Aircraft
RPAS-STE	System Synthetic Task Environment
D2T2	Distributed Dynamic Team Trust
EFA	Exploratory Factor Analysis
HAT(s)	Human-Autonomy Teaming; Human-Autonomy Team; Human-Autonomy Teams
ICM	Interaction Communication-based Measures
ITC	Interactive Team Cognition
LIWC	Linguistic Inquiry and Word Count
PRR	Perceived Relational Risk
PSR	Perceived Situational Risk
RPA	Remotely Piloted Aircraft
RPAS	Remotely Piloted Aerial System
WoZ	Wizard of Oz Methodology

SUMMARY

As artificial intelligence capabilities have improved, humans have begun teaming with autonomous agents that have the capability to communicate and make intelligent decisions that are adaptable to task situations. Trust is a core component of human-human and human-autonomy team (HAT) interaction. As with all-human teams, the amount of trust held within a HAT will impact the team's ability to perform effectively and achieve its goals. A recent theoretical framework, distributed dynamic team trust (D2T2; Huang et al., 2021), relates trust, team interaction measures, and team performance in HATs and has called for interaction-based measures of trust that go beyond traditional questionnaire-based approaches to measure the dynamics of trust in real-time. In this research, these relationships are examined by investigating HAT interaction communication-based measures (ICM; amount, frequency, affect, and “pushing” vs. “pulling” of information between team members) as a mechanism for D2T2 and tested for predictive validity using questionnaire-based trust measures as well as team performance in a three-team member remotely-piloted aerial system (RPAS) HAT synthetic task. Results suggest that ICM can be used as a measure for team performance in real-time. Specifically, ICM was a significant predictor of team performance and not trust, and trust was not a significant predictor of team performance. Exploratory factor analyses of the trust questionnaire items discovered clear differences in how human teammates characterize trust in all-human teams and HATs. Specifically for HATs, interpersonal and technical factors associated with trust in autonomous agents were found as a result of dynamic exposure to the autonomous agent by distinct stakeholders through communication. These findings confirmed the underlying

theory behind D2T2 as a framework that includes both interpersonal and technical factors related to trust in HAT along a dynamic timeline among different types of stakeholders. The findings provide some insight into the dynamic nature of trust, but continued research to discover interactive measures of trust is necessary.

CHAPTER 1. INTRODUCTION

Teamwork is traditionally defined as two or more humans working interdependently toward a common goal (Salas et al., 1992). In human-autonomy teaming (HAT), humans work interdependently with automation capable of making intelligent decisions that are adaptable to task situations (McNeese et al., 2018). HAT is becoming more prevalent in our society where humans work with automated machinery in manufacturing, smart devices at home, and autonomous agents in the military. Autonomous agents interact with human team members to achieve team-level goals and are therefore considered teammates. These agents possess the abilities to observe the environment (through some form of sensor), act upon an environment (through some form of actuator), and direct its activity toward the achievement of specific goals (Chen & Barnes, 2014). For humans and autonomous agents to work together, communication and interaction is key to achieving their goals. Furthermore, as agents become more advanced and autonomous, trust in human-autonomy teams (HATs) becomes a more pressing issue (Chen, 2018). Trust is one component in the successful deployment of autonomous systems, and the amount of trust humans hold in these autonomous agents impacts their ability to perform effectively as a team (Jian et al., 2000).

Trust is closely tied to human use and appreciation of autonomous agents or artificial systems in command-and-control systems (Sheridan, 1988). To appropriately study trust there must be some meaningful incentives at stake (i.e., betrayal or loss of something meaningful) and that the trustor and trustee must be cognizant of the risk involved (Kee & Knox, 1970). Stuck et al. (2021b) developed a model of trust with an

emphasis on how perceived risk interacts with trust. Their model is based on the definition of perceived risk by Mayer et al. (1995) that states “perceived risk involves the trustor’s belief about likelihoods of gains or losses outside of considerations that involve the relationship with the particular trustee” (p. 726). They identified two sub-types of perceived risk: perceived relational risk and perceived situational risk. Perceived relational risk (PRR) is the “belief about the probability and/or feeling that interacting with a specific system, technology, or person, with which user has a personal history or historical knowledge of, has potential negative outcomes” (p. 4). In essence PRR is the perceived risk associated with a specific system, autonomous agent, or human. Perceived situational risk (PSR) is the “belief of the probability and/or feeling that a specific task or context has potential negative outcomes based on their knowledge and experience with the task, regardless of a personal history, or historical knowledge of the system, technology, or person that may be relied on in that situation” (p. 4). In summary, PSR is the perceived risk about the negative outcomes of a task. Stuck et al. (2021a) describes how these sub-types of perceived risk can be applied in the human-robot context; however, their example is also applicable to the HAT context. In sum, a human teammate will only take a risk if their feeling of trust outweighs their perceptions of risk. The outcome of taking a risk then influences the trustor’s perceived trustworthiness in their teammate whether it be a human or autonomous agent. Results from a literature review of human-automation trust and risk reported by Stuck et al. (2021b) state that the presence of risk and participants’ PSR impacts their behavioral trust of the automation, while PRR was strongly negatively related with trust.

Trust as defined by Mayer et al. (1995) is the “willingness of a party to be vulnerable to the actions of another party based on the expectations that the other will perform a particular action important to the trustor, irrespective of the ability to monitor or control that other party” (p. 712). In these authors’ integrated model of trust, trust has the characteristics of ability, benevolence, and integrity. Ability refers to the groups of skills, competencies, and characteristics that enable a party to have influence within some specific domain (Mayer et al., 1995). Benevolence is the extent to which a trustee is believed to want to do good toward the trustor, aside from an egocentric profit motive (Mayer et al., 1995). Lastly, integrity involves the trustor’s perception that the trustee adheres to a set of principles that the trustor finds acceptable (Mayer et al., 1995). The trustor’s inherent propensity to trust will also influence the individual’s trust in the trustee prior to any interaction.

The definition of trust by Mayer et al. (1995) is primarily used for human-human trust but contains the important factors of willingness and risk that human trustors will consider when deciding to place trust in autonomous agents. As Johnson-George and Swap (1982) stated, “one of the few characteristics common to all trust situations is the willingness to take risks” (p. 1306). A definition of trust that is more applicable to HAT is Lee and See’s (2004) definition, which describes trust as “the attitude that an agent will help achieve an individual’s goals in a situation characterized by uncertainty and vulnerability” (p. 54). Lee and See (2004) view trust along a dimension of attributional abstraction varying from demonstrations of competence to the intentions of the agent. Similar to Mayer et al. (1995), Lee and See (2004) propose performance, process, and purpose as three characteristics of trust. Specifically, performance refers to the current and

historical operation of the autonomous agent, specifically the competency or expertise demonstrated by its ability to achieve goal(s) of the HAT. Performance information describes what the autonomous agent does, including characteristics such as reliability, predictability, and ability.

Marsh and Dibben (2003) identified three different layers of trust: dispositional trust, situational trust, and learned trust; as well as three sources of variability in human-autonomation trust: the human operator, the environment, and the automated system (Hoff & Bashir, 2015). Dispositional trust represents an individual's overall tendency to trust autonomous agents independent of context or specific system. It is a long-term tendency arising from both biological and environmental influences that is relatively stable over time. An individual's dispositional trust is set before any interaction with an autonomous agent and can alter or form their tendency to trust the agent. It can also vary in individuals based on interpersonal characteristics such as culture, age, gender, and personality (Hoff & Bashir, 2015). Measuring an individual's dispositional trust prior to any interaction will be necessary to capture whether an individual arrives highly or scarcely trustful toward any autonomous agent in question. An experiment by Biros, Fields, and Gunsch (2003) showed that an individual's dispositional trust in computers would predict their trust in information presented to them by an unmanned combat aerial vehicle. Results indicated that if an individual had high dispositional trust in computers, they would display more trust in the information presented to them by an unmanned aerial vehicle.

Situational trust is influenced by the environment and context-dependent variations in an individual's mental state (Hoff & Bashir, 2015). The variable external actors in the environment that can influence situational trust in autonomous agents are the type of

system, system complexity, task difficulty, workload, perceived risks, perceived benefits, organizational setting, and framing of the task. Variable internal factors such as self-confidence, subject matter expertise, mood, and attentional capacity influence the situational trust in autonomous agents. All these factors determine the degree of influence that situational trust has on interactions between an individual and an autonomous agent.

Learned trust is layer of trust formed by all the past experiences an individual had relevant to the specific autonomous agent in question (Hoff & Bashir, 2015). Learned trust is essentially the evaluation of an individual's interaction with an autonomous agent. This layer of trust is dynamic and fluctuates over time in the forms of initial learned trust, dynamic learned trust, and overall learned trust. Before interacting with a specific autonomous agent any preexisting knowledge or previous interaction(s) with the agent will bias the agent's reputation and impact an individual's initial learned trust. In several studies pointed out by Hoff and Bashir (2015), individuals displayed a tendency to trust automation more when it was portrayed as a reputable or an expert system. If any information or opinions regarding the autonomous agent were provided to an individual before interacting with it (e.g., during an informative training period), then this information would influence the initial learned trust in the agent. However, if an individual were provided an opportunity to interact with the autonomous agent after receiving the initial information (e.g., during a hands-on training mission), then the performance of the agent would impact the individual's dynamic learned trust if the evaluations are occurring during the interaction and not afterward. If the evaluation took place after the interaction, the individual would be impacting the overall learned trust in the agent. This evaluation would contain the previous evaluations from initial and dynamic learned trust making up the individual's

overall learned trust in the agent. This overall learned trust evaluation would then become the initial learned trust of the autonomous agent before the next interaction.

Trust can also be transitive in the sense that one individual's trust in an autonomous agent can be transferred to another individual. Huang et al. (2021) propose a framework that states that trust transitivity observed in human-human trust can also apply to HATs. Trust transitivity distributed throughout a team is described as interpersonal trust among all related stakeholders, including autonomous agents that can be transmitted across groups and individuals through daily conversations, newsletters and policies, and training procedures (Huang et al., 2021). In other words, the trust among team members can change through direct interactions with autonomous agents or indirectly through other human member's influence. Trust transitivity can help explain the transfer of situational and dynamic learned trust in autonomous agents from one human team member to another during interactions that are relevant to HAT operations in the current study. The trust transferred during these interactions should be reflected in the overall learned trust in the autonomous agent as the human team members evaluate their trust in the agent once the interaction is over. Initial learned trust can also be influenced by trust's transitive properties if any initial information regarding the autonomous agent given to an individual were biased by another human before the trustor were to interact with the agent in question.

Huang et al.'s (2021) proposed framework, distributed dynamic team trust (D2T2), theoretically relates trust, team interaction measures, and team performance in HATs. The framework posits that trust in autonomous agents as distributed among all related stakeholders, where interpersonal trust and human-agent trust mutually influence each other. Interaction-based team measures using behavioral and communication data (volume,

frequency, pitch, content, speech act, and flow) measured over time are hypothesized to capture D2T2 (Huang et al., 2021). The relationship between interaction-based measures and team performance in all-human teams and HATs was also explored by O'Neill et al. (2020). These authors found that human-human teams routinely outperformed HATs in part because of more efficient information sharing. This further emphasized the theory that the quality of information exchanged and communication may be important considerations for HATs and their future performance (O'Neill et al., 2020).

Trust was shown to be related to the performance of HATs in an experiment by McNeese et al. (2019). The experiment sought to understand trust and its relations to team performance using a Wizard of Oz (WoZ) methodology to simulate an autonomous agent as a team member in a remotely piloted aircraft system environment (RPAS; Kelley, 1983, December 13—15; McNeese et al., 2019). The WoZ methodology places an experimenter in the role of an autonomous agent teammate while ensuring that the participants believe that they are working with an authentic autonomous agent. Using WoZ, research can be conducted by following a script instead of programming an autonomous agent, allowing for controlled behavior that might be beyond the technical capabilities of a programmed agent. In the McNeese et al. (2019) study, the HAT was comprised of a “synthetic teammate” (WoZ) and two human teammates who had to communicate with one another to take reconnaissance photos of enemy targets over a series of 40-minute missions. Their results showed that lower performing teams had lower levels of trust in the “synthetic teammate”. However, it was unclear whether lower team performance predicted lower levels of trust or if the HATs with lower levels of trust predicted lower team performance.

1.1 The Current Study

In a study by Lee and Kolodge (2020) text-based analyses of conversations surrounding humans' trust in autonomous vehicles showed that communication can be a way to unobtrusively measure trust between humans and autonomy. Communication is a directly observable measure of team cognition that has been consistently tied to team performance (Cooke et al., 2013), which does not suffer from the subjectivity of most trust measures. If communication can be tied to trust, then it would provide a more objective measure with potential for real-time analysis, as theorized by the D2T2 framework (Huang et al., 2021). In the current study, it is hypothesized that team communication can be objectively tied to subjective trust measures while also predicting team performance in a simulated RPAS HAT, in which two human operators (navigator; photographer) work with either an autonomous agent or a trained experimenter playing the role of the pilot over a series of aerial reconnaissance missions.

The current study does not focus on the manipulation of training with an autonomous agent vs. human experimenter pilot, which was the original aim of the study. Rather, the current study focuses on the relations among measured variables to tie interaction-based metrics to trust and team performance. Specifically, exploratory factor analysis of responses to trust questionnaire items, followed by a regression of interaction communication-based measures (ICM) on the resulting trust scores and an objective measure of team performance are analyzed. Thus, the goal of the current study is to determine if we can predict both trust and team performance using objective communication measures of amount, frequency, affect, and “pushing” vs. “pulling” of information between team members over time, as predicted by the D2T2 framework.

1.2 Hypotheses

This study aims to test the following predictions related to trust, team performance, and ICM in HATs.

1.2.1 Hypothesis 1

Objective ICM will predict both subjective human-human trust measures and an objective team performance measure in the all-human condition.

1.2.2 Hypothesis 2

Objective ICM will predict both subjective trust in autonomy and an objective measure of team performance in the human-autonomy condition.

CHAPTER 2. METHOD

2.1 Participants

Twenty-one dyads comprised of 42 participants were recruited from Georgia Institute of Technology and its surrounding area. These dyads teamed with either an autonomous agent or trained experimenter to form three-member teams. All teams participated in one six-hour session consisting of training and four 40-minute missions. Participants had normal or corrected-to-normal vision and were required to be fluent in English. Ages ranged from 18 to 31 years ($M = 20.5$, $SD = 2.9$) across 21 males, 20 females, and one non-binary person. Each participant was compensated with a combination of \$10.00 per hour or 1 hour of research credit per hour of participation.

2.2 Materials

The experiment was conducted in the Cognitive Engineering Research on Team Tasks Remote Piloted Aircraft System Synthetic Task Environment (CERTT-RPAS-STE; Cooke & Shope, 2005) located at Georgia Tech. The CERTT-RPAS-STE is comprised of three task-role stations and four experimenter stations. The objective is to take photographs of color-coded strategic target waypoints while avoiding color-coded hazard waypoints over a series of 40-minute missions. Team performance (0-1000) is scored based upon number of successful target photos, resource (fuel; film), usage, and penalty points deducted if they encounter a hazard, warning, or alarm.

The first role, pilot (AVO) controls and monitors the altitude and airspeed of the RPA, vehicle heading, fuel, gears, and flaps, and interacts with the photographer to

negotiate altitude and airspeed to take a clear picture of the various target waypoints. This role was played by either an autonomous agent (“synthetic” teammate) or a human experimenter. The synthetic teammate was developed using the ACT-R cognitive modeling architecture to simulate human cognition and interacts with the human teammates in the CERTT-RPAS-STE using text chat (Ball et al., 2010). The synthetic teammate is capable of deciding its own course of action based on its experiences during the dynamic task situation and is responsible for all aspects of the role (McNeese et al., 2018). The synthetic teammate was not developed with explicit teamwork skills, yet it is a critical part of the team and cannot be set aside if the team expects to perform well (McNeese et al., 2018). For teams in the synthetic teammate condition, during the last mission of the experiment, the pilot role was assumed by a trained experimenter; however, given the motivation of the current study to validate ICM as predictors of trust and team performance, this manipulation was not directly evaluated in the current study. The participants were aware of when they were working with the synthetic teammate or the human pilot.

The second role, navigator (DEMPC), creates a dynamic flight plan and notifies the pilot of information regarding waypoints, including waypoint name, altitude restrictions, airspeed restrictions, and effective target radius. The third role, photographer (PLO), monitors and adjusts camera settings to take target photos and sends feedback to the other teammates regarding photo quality. These two roles were occupied by the participants. All team members communicated using a text-chat interface. One experimenter played the role of intelligence, who communicated with the team if they asked for help. The remaining

two experimenter stations were used to log information within the task environment that is beyond the scope of the current study.

2.3 Procedure

Before arriving, each team was randomly assigned to an experimental condition (synthetic teammate pilot vs. trained experimenter pilot). After providing informed consent, participants were instructed to fill out the first set of trust questionnaires. Participant training consisted of an individual 30-minute interactive PowerPoint training module focusing on each participant's role, followed by a 30-minute hands-on team training mission to familiarize themselves with the CERTT-RPAS-STE. Experimenters coached the participants while following a script to ensure each participant understood how to communicate, their roles, and the task. Teams then engaged in Missions 1 and 2 followed by a short break. After the break, participants performed Mission 3 and then filled out the second set of trust questionnaires. For Mission 4, the pilot role was always assumed by an experimenter to examine transfer from synthetic pilot to human pilot. However, this manipulation is not directly examined in the current study. The last set of trust questionnaires were then completed. Participants were then debriefed and paid or given credit for their participation in the 6-hour study.

Table 1 – Experimental Session

	Human-Autonomy Condition	All-Human Condition
	Pilot Role	
Questionnaire Session 1		
Training	Synthetic Teammate	Experimenter
Mission 1	Synthetic Teammate	Experimenter
Mission 2	Synthetic Teammate	Experimenter

Table 1 continued

Mission 3	Synthetic Teammate	Experimenter
Questionnaire Session 2		
Mission 4	Experimenter	Experimenter
Questionnaire Session 3		

2.4 Measures

2.4.1 ICM

The ICM included the amount, frequency, affect, and “pushing” vs. “pulling” of chat message information. Besides affect, all of these component measures were collected during each Mission from the messages within the chat log embedded in the CERTT-RPAS-STE. “Pushing” message information refers to team verbal behaviors related to sending information to other team members whereas “pulling” team verbal behaviors are related to asking for information (McNeese et al., 2018). The team verbal behaviors listed in Table 2 were tagged by two experimenters resulting in a numerical amount of “pushing” and “pulling” team verbal behaviors in the CERTT-RPAS-STE. Inter-rater reliability for these “push” and “pull” behaviors are listed in Table 3. The measures of message amount, frequency, “pushing” and “pulling” were each aggregated to the team level (i.e., summed across missions and then divided by the number of missions when brought to the team level). Message affect was analyzed using Linguistic Inquiry and Word Count (LIWC; Boyd et al., 2022), wherein LIWC outputs a number based on positive and negative tone. Numbers above 50 suggest a more positive emotional tone whereas numbers below 50 suggest a negative emotional tone (Boyd et al., 2022). Affect was also aggregated to the team level. ICM is a total sum of message amount, frequency, affect, and “pushing” and “pulling” of message information at the team level. ICM and its component measures were

brought to the team level to ensure that all measures including team trust and team performance were at the same level of analysis.

Table 2 – Team Verbal Behaviors

Behaviors	Push/Pull	Description
General Status Update	Push	Informing other team members about current status
Suggestions	Push	Making suggestions to the other team members
Planning Ahead	Push	Anticipating next steps and creating rules for future encounters
Repeated Request	Pull	Requesting the same information or action from other team member(s)
Inquiry About Status of Others	Pull	Inquiring about current status of others and expressing concerns
<i>Note.</i> This table is a modified table from McNeese et al. (2018).		

Table 3 – Team Verbal Behavior Inter-Rater Reliability Cohen’s K

Behaviors	κ
General Status Update	0.682
Suggestions	0.671
Planning Ahead	0.474
Repeated Request	0.666
Inquiry About Status of Others	0.760

2.4.2 Team Trust

Team trust is an aggregated score of trust from both the navigator and photographer roles at the team level. Two trust questionnaires are analyzed for the purposes of this study, once before the training session and again after Mission 3. The third set of trust questionnaires were not analyzed for this study because the autonomous agent pilot was replaced by a trained human experimenter in the fourth Mission. The first questionnaire was a modified trust questionnaire originally developed by Mayer and Gavin (2005; Appendix A). This questionnaire was modified by Demir et al. (2021) to fit the HAT context and consists of 25 items regarding trust towards either human or autonomous teammates with a Likert scale ranging from “1” = Strongly Agree to “5” = Strongly Disagree. The second questionnaire was the Checklist for Trust between People and Automation Scale (Appendix B; Jian et al., 2000). This questionnaire consists of 12 items with a Likert scale ranging from “1” = Not at All to “7” = Extremely. To keep the Likert scale ratings in the same direction (“1” = Strongly Disagree or Not at All), the modified questionnaire by Mayer and Gavin (2005) was reverse scored.

There were missing data for some items of the team trust measure from Team 3, Team 11, and Team 16. A total of 39 out of 3108 items were mean replaced. To correct for missing data, all missing items were averaged across role per condition and session of that specific item. For example, if there were missing data for the 11th item on the Mayer and Gavin (2005) questionnaire for the photographer on Team 3 (human-autonomy condition) at the second session, then the scores for all photographers in the human-autonomy condition at the second session for the 11th item were averaged to replace the missing item score. This mean imputation was calculated accordingly to account for the variance across

items in the exploratory factor analysis and the later aggregation of trust scores across questionnaire sessions to the team level.

2.4.3 Team Performance

Team performance is an objective outcome measure scored out of 1,000 per Mission and is scored at the team level. At each Mission, teams begin with 1,000 points, and points are deducted based on a weighted composite of team level parameters, including the number of missed targets, rate of good photos taken per minute, film and fuel resource consumption, and time spent in warning and alarm states.

2.5 Design

Although the design of the experiment was motivated by the question of whether HAT task acquisition transfers to all-human team performance, the purpose of the current study is to factor analyze responses from trust questionnaires and regress ICM on the resulting trust scores and an objective measure of team performance to determine if we can predict subjective measures of trust using objective measures of team communication. In this experiment, each team completed one training and four experimental missions and answered three sets of trust questionnaires (Table 1). However, for the current study only the measures and responses from the first three missions and two sets of trust questionnaires were used. First, two separate exploratory factor analyses (EFA) were conducted on the 37 trust items from the questionnaires to reveal the underlying factor structure or factor scores in the human-autonomy and all-human condition. Then, two third variable models each containing three regression models were tested to establish ICM as a

third variable that explains the relationship between team trust and team performance in HATs and all-human teams.

2.5.1 *The Third Variable Problem*

There are three fundamental types of third variable problems: common cause, mediation, and moderation. In common cause models, two variables are related due to their separate relationships to the same third variable. In mediation models, the effect of an independent variable on a dependent variable “goes through” a third variable, and in moderation models the relationship between two variables is conditioned on the value of a third variable (Baron & Kenny, 1986). In partially redundant common cause models, the two variables are still related to one another. However, the models tested posit that ICM will be a third variable that explains the relationship between team trust and team performance where team trust is predictive of team performance. Contrary to partially redundant common cause models, the hypothesized model specifies a relationship between the two variables in question (i.e., team trust and team performance). Figure 1 shows the third variable model that will be tested in the all-human and human-autonomy conditions.

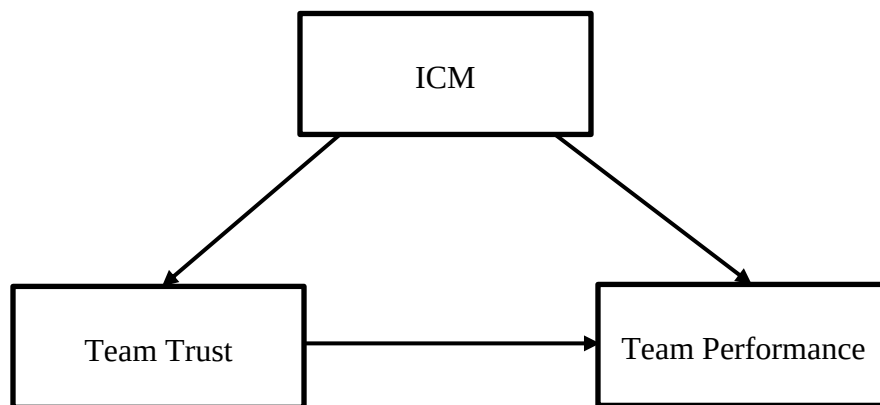


Figure 1 – Third variable model hypothesized for both conditions.

Following Baron and Kenny's (1986) method for testing third variable models with path diagrams, three regression models were tested for the proposed third variable model. The first model used in Step 1 is pictured in Figure 2, the second model used in Step 2 is pictured in Figure 3, and the third model used in Step 3 is pictured in Figure 4. In Step 1, the model is tested where team trust is the independent variable and team performance is the dependent variable. In Step 2, the model is tested where ICM is the independent variable and team trust is the dependent variable. In Step 3, the model is tested where ICM and team trust are independent variables and team performance is the dependent variable. To establish ICM as a third variable that could be used in place of team trust, not only should the first and second model be significant, but in the third model, the second model (ICM→Performance) should continue to be significant whereas the first model (Trust→Performance) should no longer be significant. Objective ICM could then be used in place of subjective trust questionnaires to account for the underlying factors found in each EFA and to predict team performance.

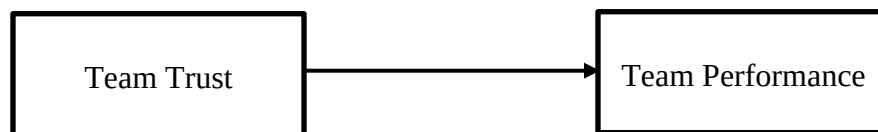


Figure 2 – Model 1 for both conditions.

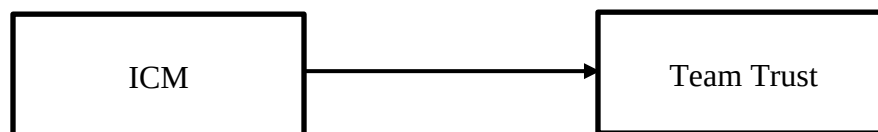


Figure 3 – Model 2 for both conditions.

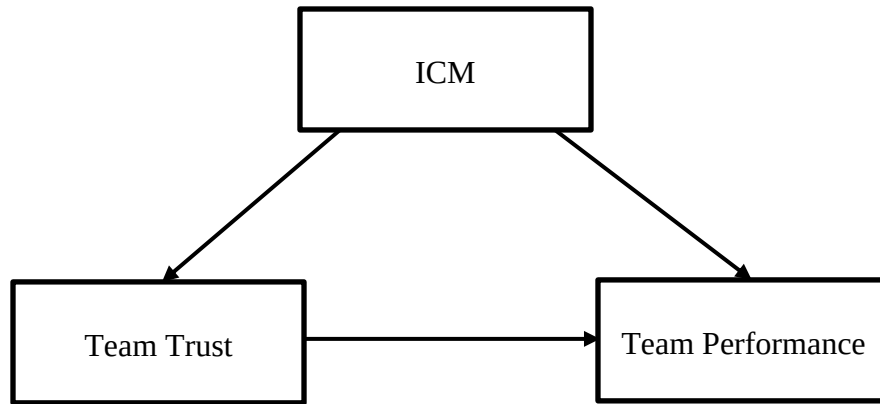


Figure 4 – Model 3 for both conditions.

CHAPTER 3. RESULTS

3.1 Exploratory Factor Analyses

To uncover the underlying factor structure of the 37 trust items, responses from both questionnaire sessions were used in both the analysis of the HATs and all-human teams. The EFAs were based on principle axis factoring with Kaiser's varimax rotation to reduce dimensionality and derive an appropriate number of factors.

3.1.1 Exploratory Factor Analysis in All-Human Condition

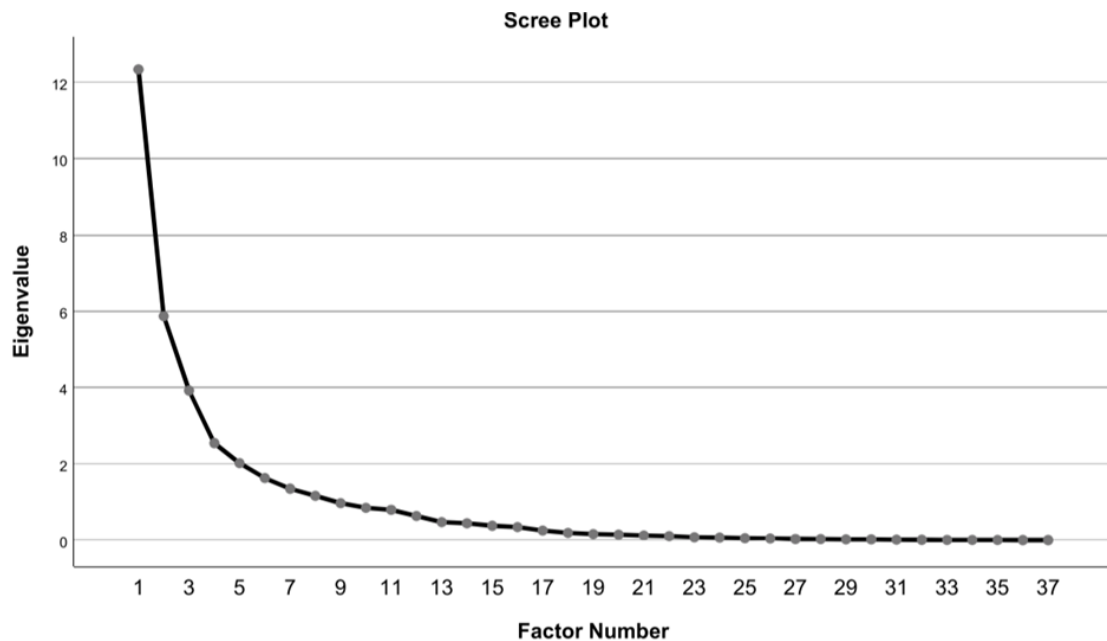


Figure 5 – All-human scree plot.

According to the EFA findings for the all-human teams, 83.30% of the total variance was accounted for by eight factors. The all-human scree plot (Figure 5) shows a notable drop at the fourth factor, which indicates that the three factors above explain most of the variance. Therefore, the first three factors were retained, which cumulatively

accounted for 59.81% of the variance and 33.35%, 15.86%, and 10.60% respectively. As shown in Table 4, Factor 1 represents the *trust in the CERTT-RPAS-STE system (without synthetic teammate)*, Factor 2 represents the *trust in (human) teammates*, and Factor 3 represents the *human teammates' desire to monitor respective teammates*.

Table 4 – Factor Loadings on the First Three Factors in the All-Human Condition

Factor	Item	Factor Loading
Trust in the CERTT-RPAS-STE System (Without Synthetic Teammate)	I can trust the system	.930
	The system is dependable	.917
	I am confident in the system	.913
	The system has integrity	.889
	The system provides security	.883
	The system is reliable	.881
	I am suspicious of the system's intent, action, or outputs	.778
	I am wary of the system	.770
	The system's actions will have a harmful or injurious outcome	.762
	The system is deceptive	.760
	The system behaves in an underhanded manner	.741
	I am familiar with the system	.392
Trust in (Human) Teammates	I felt the AVO was reliable	.869
	If the PLO/DEMPC asked why a problem happened, I would speak freely even if I were partly to blame	.827
	I trusted the AVO	.818
	If the AVO asked why a problem happened, I would speak freely even if I were partly to blame	.805
	I would tell the PLO/DEMPC about mistakes I have made on the team task, even if they could damage my reputation	.804
	I enjoyed working with the AVO	.800
	I would tell the AVO about mistakes I have made on the team task, even if they could damage my reputation	.786
	I would be comfortable giving AVO a task or problem which was critical to me, even if I could not monitor his/her/its actions	.770
	I would share my opinion about sensitive issues with the AVO even if my opinion were unpopular	.732
	I would share my opinion about sensitive issues with the PLO/DEMPC even if my opinion were unpopular	.728

Table 4 continued

Human Teammates' Desire to Monitor Respective Teammates	I would be comfortable giving PLO/DEMPC a task or problem which was critical to me, even if I could not monitor his/her/its actions	.683
	While chatting with AVO, it felt like I was talking to a real person	.608
	I really wish I had a good way to keep an eye on the PLO/DEMPC	.698
	I really wish I had a good way to keep an eye on the AVO	.666

3.1.2 Exploratory Factor Analysis in Human-Autonomy Condition

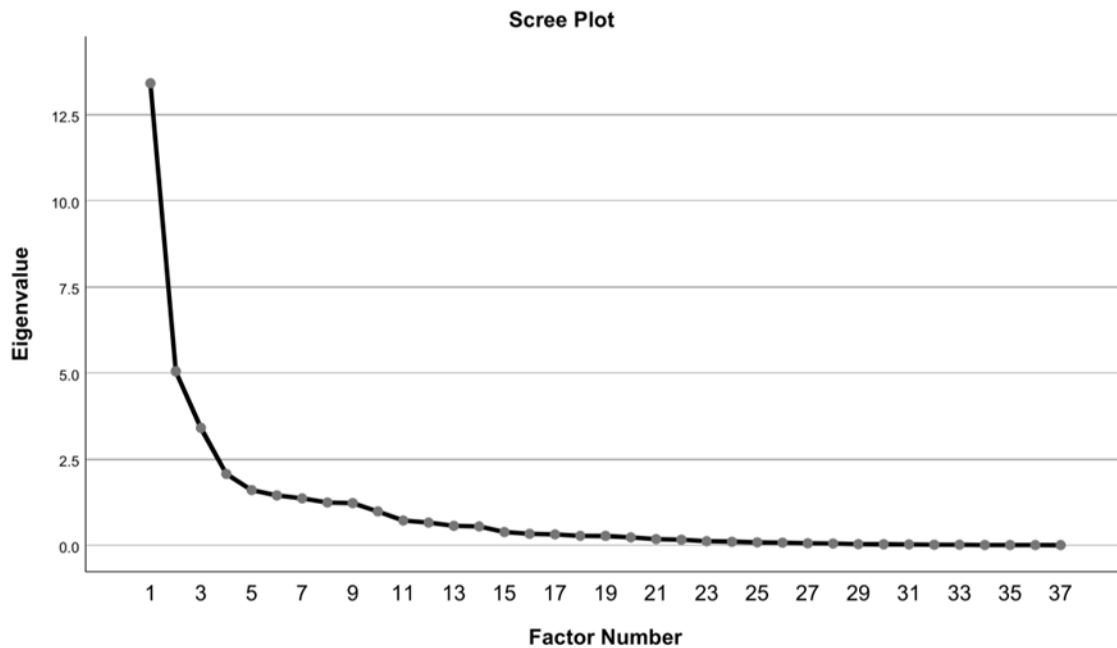


Figure 6 – HAT scree plot.

According to the EFA findings for the HATs, 83.29% of the total variance was accounted for by nine factors. The scree plot (Figure 5) shows a notable drop at the fourth factor, which indicates that the three factors above explain most of the variance. Therefore,

the first three factors were retained, which cumulatively accounted for 59.11% of the variance and 36.25%, 13.65%, and 9.21% respectively. As shown in Table 4, Factor 1 represents the *trust in the synthetic teammate*, Factor 2 represents the *human teammates' openness to admit mistakes*, and Factor 3 represents *popularity and reputation among teammates*.

Table 5 – Factor Loadings on the First Three Factors in the Human-Autonomy Condition

Factor	Item	Factor Loading
Trust in the Autonomous Teammate	I can trust the system	.931
	I felt the AVO was reliable	.926
	The system is reliable	.918
	I am confident in the system	.907
	The system is dependable	.865
	I enjoyed working with the AVO	.851
	The system provides security	.850
	I trusted the AVO	.838
	While chatting with AVO, it felt like I was talking to a real person	.819
	I am wary of the system	.770
	The system's actions will have a harmful or injurious outcome	.766
	If I had my way, I would not let the AVO have any influence over issues that are important to me	.716
	If someone questioned the AVOs motives, I would give the AVO the benefit of the doubt	.690
	The system has integrity	.642
	I am suspicious of the system's intent, action, or outputs	.616
	I really wish I had a good way to keep an eye on the AVO	.600
	I would be comfortable giving AVO a task or problem which was critical to me, even if I could not monitor his/her/its actions	.567
Human Teammates' Openness to Admit Mistakes	If the AVO asked why a problem happened, I would speak freely even if I were partly to blame	.894
	If the PLO/DEMPC asked why a problem happened, I would speak freely even if I were partly to blame	.702
	I would tell the AVO about mistakes I have made on the team task, even if they could damage my reputation	.499

Table 5 continued

Popularity and Reputation Among Teammates	I felt the AVO displayed feminine qualities	-.482
	I would share my opinion about sensitive issues with the AVO even if my opinion were unpopular	.818
	I would share my opinion about sensitive issues with the PLO/DEMPC even if my opinion were unpopular	.765
	I would tell the PLO/DEMPC about mistakes I have made on the team task, even if they could damage my reputation	.653

3.2 Third Variable Models

To address the research question of whether ICM is a third variable that can be used in place of team trust and predict team performance, three regression models were tested for both the all-human and human-autonomy condition.

3.2.1 ICM, Team Trust, and Team Performance in All-Human Teams

Table 6 – Summary of Regression Analysis for Model 1 in the All-Human Condition

	<i>B</i>	<i>SE B</i>	<i>β</i>	<i>t</i>	<i>p</i>
Team Trust	-0.371	1.772	-0.562	-0.209	.840
<i>Note.</i> This model tests if team trust predicts team performance in all-human teams.					

The results from the first regression model (Table 6) indicate that team trust did not significantly explain a proportion of variance in team performance, $R^2 = 0.005$, $F(1, 8) = 0.04$, $p = .840$. The results from the second regression model (Table 7) indicate that ICM

Table 7 – Summary of Regression Analysis for Model 2 in the All-Human Condition

	<i>B</i>	<i>SE B</i>	<i>β</i>	<i>t</i>	<i>p</i>
ICM	-0.108	0.095	-0.372	-1.135	.289
<i>Note.</i> This model tests if ICM predicts team trust in all-human teams.					

did not significantly explain a proportion of variance in team trust, $R^2 = 0.14$, $F(1, 8) = 1.29$, $p = .289$. The third regression model (Table 8) where ICM and team trust were

predictors of team performance did not significantly explain a proportion of variance in team performance, $R^2 = 0.47$, $F(2, 7) = 3.13$, $p = .107$. Yet, ICM was a significant predictor of team performance, $\beta = 0.74$, $t(7) = 2.49$, $p < .05$, which indicates that, on average, each per unit of ICM is associated with a 0.74 *SD* increase in team performance. Team trust was not a significant predictor in this model, $\beta = 0.20$, $t(7) = 0.68$, $p = .520$.

Table 8 – Summary of Regression Analysis for Model 3 in the All-Human Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
ICM	1.076	0.433	0.736	2.488	.042
Team Trust	1.006	1.487	0.200	0.677	.520
<i>Note.</i> This model tests if ICM and team trust predicts team performance in all-human teams.					

3.2.1.1 Post-Hoc Analysis: Third Variable Analysis Using the Team Trust Scores from Session 1 in All-Human Teams

Table 9 – Summary of Regression Analysis for Model 1 using Team Trust Scores from Session 1 in the All-Human Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Team Trust	-0.851	1.227	-0.238	-0.693	.508
<i>Note.</i> This model tests if team trust taken from questionnaire session 1 predicts team performance in all-human teams.					

To further examine the relationship between ICM, team trust, and team performance, the team trust scores from the first questionnaire session were isolated and implemented in a follow-up third variable analysis. The first questionnaire session took place before the participants began training on the CERTT-RPAS-STE. The results from the first regression model (Table 9) indicate that team trust in Session 1 did not significantly explain a proportion of variance in team performance, $R^2 = 0.06$, $F(1, 8) = 0.48$, $p = .508$.

Table 10 – Summary of Regression Analysis for Model 2 using Team Trust Scores from Session 1 in the All-Human Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
ICM	-0.169	0.132	-0.412	-1.278	.237
<i>Note.</i> This model tests if ICM predicts team trust taken from questionnaire session 1 in all-human teams.					

The results from the second regression model (Table 10) indicate that ICM did not significantly explain a proportion of variance in team trust, $R^2 = 0.17$, $F(1, 8) = 1.63$, $p = .237$. The third regression model (Table 11) where ICM and team trust were predictors of team performance did not significantly explain a proportion of variance in team performance, $R^2 = 0.44$, $F(2, 7) = 2.74$, $p = .132$. ICM was not a significant predictor of team performance, $\beta = 0.68$, $t(7) = 2.18$, $p = .065$. Likewise, team trust was not a significant predictor in this model, $\beta = 0.04$, $t(7) = 0.13$, $p = .898$.

Table 11 – Summary of Regression Analysis for Model 3 using Team Trust Scores from Session 1 in the All-Human Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
ICM	0.992	0.454	0.678	2.184	.065
Team Trust	0.147	1.110	0.041	0.133	.898
<i>Note.</i> This model tests if ICM and team trust taken from questionnaire session 1 predicts team performance in all-human teams.					

3.2.1.2 Post-Hoc Analysis: Third Variable Analysis Using the Team Trust Scores from Session 2 in All-Human Teams

Table 12 – Summary of Regression Analysis for Model 1 using Team Trust Scores from Session 2 in the All-Human Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Team Trust	1.039	1.830	0.197	0.568	.586
<i>Note.</i> This model tests if team trust taken from questionnaire session 2 predicts team performance in all-human teams.					

To further examine the relationship between ICM, team trust, and team performance, the team trust scores from the second questionnaire session were isolated and implemented in a follow-up third variable analysis. The second questionnaire session took

place after the participants finished Mission 3. The results from the first regression model (Table 12) indicate that team trust in Session 2 did not significantly explain a proportion of variance in team performance, $R^2 = 0.04$, $F(1, 8) = 0.32$, $p = .586$. The results from the

Table 13 – Summary of Regression Analysis for Model 2 using Team Trust Scores from Session 2 in the All-Human Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
ICM	-0.048	0.096	-0.174	-0.500	.630
<i>Note.</i> This model tests if ICM predicts team trust taken from questionnaire session 2 in all-human teams.					

second regression model (Table 13) indicate that ICM did not significantly explain a proportion of variance in team trust, $R^2 = 0.03$, $F(1, 8) = 0.25$, $p = .630$. The third regression model (Table 14) where ICM and team trust were predictors of team performance did not significantly explain a proportion of variance in team performance, $R^2 = 0.54$, $F(2, 7) = 4.08$, $p = .067$. Yet, ICM was a significant predictor of team performance, $\beta = 0.72$, $t(7) = 2.75$, $p < 0.05$, which indicates that, on average, each per unit of ICM is associated with a 0.72 *SD* increase in team performance. Team trust was not a significant predictor in this model, $\beta = 0.32$, $t(7) = 1.23$, $p = .257$.

Table 14 – Summary of Regression Analysis for Model 3 using Team Trust Scores from Session 2 in the All-Human Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
ICM	1.049	0.381	0.718	2.750	.029
Team Trust	1.700	1.377	0.322	1.234	.257
<i>Note.</i> This model tests if ICM and team trust taken from questionnaire session 2 predicts team performance in all-human teams.					

3.2.1.3 Post-Hoc Analysis: Regression Analyses of the Average Amount of Messages

Component of ICM as it Predicts Team Trust and Team Performance in All-

Human Teams

Table 15 – Summary of Regression Analysis for Average Amount of Messages Predicting Team Trust in the All-Human Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Average Amount of Messages	-0.483	0.234	-0.590	-2.067	.073
<i>Note.</i> This model tests if the ICM component, average amount of messages, predicts team trust in all-human teams.					

To examine which ICM component measure was the most useful for predicting team trust and team performance in all-human teams, the average amount of messages per team was isolated and tested. The results from the regression model where the average amount of messages predicts team trust (Table 15) indicate that the average amount of messages component does not significantly explain a proportion of variance in team trust, $R^2 = 0.35$, $F(1, 8) = 4.27$, $p = .073$. Similarly, the results from the regression model where the average amount of messages predicts team performance (Table 16) also indicate that the average amount of messages component does not significantly explain a proportion of variance in team performance, $R^2 = 0.29$, $F(1, 8) = 3.34$, $p = .105$.

Table 16 – Summary of Regression Analysis for Average Amount of Messages Predicting Team Performance in the All-Human Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Average Amount of Messages	2.235	1.222	0.543	1.829	.105
<i>Note.</i> This model tests if the ICM component, average amount of messages, predicts team performance in all-human teams.					

3.2.1.4 Post-Hoc Analysis: Regression Analyses of the Team Message Frequency

Component of ICM as it Predicts Team Trust and Team Performance in All-Human Teams

Table 17 – Summary of Regression Analysis for Team Message Frequency Predicting Team Trust in the All-Human Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Team Message Frequency	-829.460	562.210	-0.462	-1.475	.178

Table 17 continued

<i>Note.</i> This model tests if the ICM component, team message frequency, predicts team trust in all-human teams.

To examine which ICM component measure was the most useful for predicting team trust and team performance in all-human teams, the team message frequency was isolated and tested. The results from the regression model where the team message frequency predicts team trust (Table 17) indicate that team message frequency does not significantly explain a proportion of variance in team trust, $R^2 = 0.21$, $F(1, 8) = 2.18$, $p = .178$. The results from the regression model where team message frequency predicts team performance (Table 18) indicate that team message frequency significantly explains a proportion of variance in team performance, $R^2 = 0.41$, $F(1, 8) = 5.61$, $p < .05$. This component of ICM was a significant predictor of team performance, $\beta = 0.64$, $t(8) = 2.37$, $p < .05$, which indicates that, on average, each message per second is associated with a 0.64 *SD* increase in team performance.

Table 18 – Summary of Regression Analysis for Team Message Frequency Predicting Team Performance in the All-Human Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Team Message Frequency	5783.000	2442.900	0.642	2.368	.045
<i>Note.</i> This model tests if the ICM component, team message frequency, predicts team performance in all-human teams.					

3.2.1.5 Post-Hoc Analysis: Regression Analyses of the Team Affect Component of ICM as it Predicts Team Trust and Team Performance in All-Human Teams

Table 19 – Summary of Regression Analysis for Team Affect Predicting Team Trust in the All-Human Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Team Affect	-0.277	0.444	-0.215	-0.624	.550
<i>Note.</i> This model tests if the ICM component, team affect, predicts team trust in all-human teams.					

To examine which ICM component measure was the most useful for predicting team trust and team performance in all-human teams, the team affect component was isolated and tested. The results from the regression model where team affect predicts team trust (Table 19) indicate that team affect does not significantly explain a proportion of variance in team trust, $R^2 = 0.05$, $F(1, 8) = 0.39$, $p = .550$. Similarly, the results from the regression model where team affect predicts team performance (Table 20) also indicate that team affect does not significantly explain a proportion of variance in team performance, $R^2 = 0.24$, $F(1, 8) = 2.49$, $p = .153$.

Table 20 – Summary of Regression Analysis for Team Affect Predicting Team Performance in the All-Human Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Team Affect	3.145	1.993	0.487	1.578	.153
<i>Note.</i> This model tests if the ICM component, team affect, predicts team performance in all-human teams.					

3.2.1.6 Post-Hoc Analysis: Regression Analyses of the Team Pushing Verbal Behavior Component of ICM as it Predicts Team Trust and Team Performance in All-Human Teams

Table 21 – Summary of Regression Analysis for Team Pushing Verbal Behavior Predicting Team Trust in the All-Human Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Team Pushing Verbal Behavior	-0.082	0.232	-0.123	-0.352	.734
<i>Note.</i> This model tests if the ICM component, team pushing verbal behavior, predicts team trust in all-human teams.					

To examine which ICM component measure was the most useful for predicting team trust and team performance in all-human teams, the team pushing verbal behavior component was isolated and tested. The results from the regression model where team pushing verbal behavior predicts team trust (Table 21) indicate that team pushing verbal

behavior does not significantly explain a proportion of variance in team trust, $R^2 = 0.02$, $F(1, 8) = 0.12$, $p = .734$. The results from the regression model where team pushing verbal behavior predicts team performance (Table 22) indicate that team pushing verbal behavior significantly explains a proportion of variance in team performance, $R^2 = 0.48$, $F(1, 8) = 7.30$, $p < .05$. This component of ICM was a significant predictor of team performance, $\beta = 0.69$, $t(8) = 2.70$, $p < .05$, which indicates that, on average, each message containing a team push verbal behavior is associated with a 0.69 *SD* increase in team performance.

Table 22 – Summary of Regression Analysis for Team Pushing Verbal Behavior Predicting Team Performance in the All-Human Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Team Pushing Verbal Behavior	2.296	0.850	0.691	2.701	.027
<i>Note.</i> This model tests if the ICM component, team pushing verbal behavior, predicts team performance in all-human teams.					

3.2.1.7 Post-Hoc Analysis: Regression Analyses of the Team Pulling Verbal Behavior

Component of ICM as it Predicts Team Trust and Team Performance in All-Human Teams

Table 23 – Summary of Regression Analysis for Team Pulling Verbal Behavior Predicting Team Trust in the All-Human Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Team Pulling Verbal Behavior	-0.690	0.601	-0.376	-1.148	.284
<i>Note.</i> This model tests if the ICM component, team pulling verbal behavior, predicts team trust in all-human teams.					

To examine which ICM component measure was the most useful for predicting team trust and team performance in all-human teams, the team pulling verbal behavior component was isolated and tested. The results from the regression model where team pulling verbal behavior predicts team trust (Table 23) indicate that team pulling verbal behavior does not significantly explain a proportion of variance in team trust, $R^2 = 0.14$,

$F(1, 8) = 1.318, p = .284$. Similarly, the results from the regression model where team pulling verbal behavior predicts team performance (Table 24) also indicate that team pulling verbal behavior does not significantly explain a proportion of variance in team performance, $R^2 = 0.12, F(1, 8) = 1.066, p = .332$.

Table 24 – Summary of Regression Analysis for Team Pulling Verbal Behavior Predicting Team Performance in the All-Human Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Team Pulling Verbal Behavior	3.162	3.062	0.343	1.033	.332
<i>Note.</i> This model tests if the ICM component, team pulling verbal behavior, predicts team performance in all-human teams.					

3.2.2 ICM, Team Trust, and Team Performance in HATs

Table 25 – Summary of Regression Analysis for Model 1 in the Human-Autonomy Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Team Trust	0.822	1.829	0.148	0.449	.664
<i>Note.</i> This model tests if team trust predicts team performance in HATs.					

The results from the first regression model (Table 25) indicate that team trust did not significantly explain a proportion of variance in team performance, $R^2 = 0.02, F(1, 9) = 0.20, p = .644$. The results from the second regression model (Table 26) indicate that

Table 26 – Summary of Regression Analysis for Model 1 in the Human-Autonomy Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
ICM	-0.114	0.125	-0.291	-0.913	.385
<i>Note.</i> This model tests if ICM predicts team trust in HATs.					

ICM did not significantly explain a proportion of variance in team trust, $R^2 = 0.08, F(1, 9) = 0.83, p = .385$. The third regression model (Table 27) where ICM and team trust were predictors of team performance did not significantly explain a proportion of variance in

team performance, $R^2 = 0.51$, $F(2, 8) = 4.12$, $p = .059$. Yet, ICM was a significant predictor of team performance, $\beta = 0.73$, $t(8) = 2.81$, $p < .05$, which indicates that, on average, each per unit of ICM is associated with a 0.73 *SD* increase in team performance. Team trust was not a significant predictor in this model, $\beta = 0.36$, $t(8) = 1.39$, $p = .203$.

Table 27 – Summary of Regression Analysis for Model 2 in the Human-Autonomy Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
ICM	1.584	0.565	0.728	2.806	.023
Team Trust	1.998	1.440	0.360	1.388	.203
<i>Note.</i> This model tests if ICM and team trust predicts team performance in HATs.					

3.2.2.1 Post-Hoc Analysis: Third Variable Analysis Using the Team Trust Scores from Session 1 in HATs

Table 28 – Summary of Regression Analysis for Model 1 using Team Trust Scores from Session 1 in the Human-Autonomy Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Team Trust	-0.383	2.042	-0.062	-0.188	.855
<i>Note.</i> This model tests if team trust taken from questionnaire session 1 predicts team performance in HATs.					

To further examine the relationship between ICM, team trust, and team performance, the team trust scores from the first questionnaire session were isolated and implemented in a follow-up third variable analysis. The first questionnaire session took place before the participants began training on the CERTT-RPAS-STE. The results from the first regression model (Table 28) indicate that team trust in Session 1 did not significantly explain a proportion of variance in team performance, $R^2 = 0.004$, $F(1, 9) = 0.04$, $p = .855$. The results from the second regression model (Table 29) indicate that ICM

Table 29 – Summary of Regression Analysis for Model 2 using Team Trust Scores from Session 1 in the Human-Autonomy Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
ICM	-0.136	0.109	-0.385	-1.251	.243
<i>Note.</i> This model tests if ICM predicts team trust taken from questionnaire session 1 in HATs.					

did not significantly explain a proportion of variance in team trust, $R^2 = 0.15$, $F(1, 9) = 1.56$, $p = .243$. The third regression model (Table 30) where ICM and team trust were predictors of team performance did not significantly explain a proportion of variance in team performance, $R^2 = 0.43$, $F(2, 8) = 2.96$, $p = .109$. Yet, ICM was a significant predictor of team performance, $\beta = 0.70$, $t(8) = 2.42$, $p < .05$, which indicates that, on average, each per unit of ICM is associated with a 0.70 *SD* increase in team performance. However, team trust was not a significant predictor in this model, $\beta = 0.21$, $t(8) = 0.72$, $p = .494$.

Table 30 – Summary of Regression Analysis for Model 3 using Team Trust Scores from Session 1 in the Human-Autonomy Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
ICM	1.530	0.632	0.703	2.422	.042
Team Trust	1.279	1.782	0.208	0.717	.494
<i>Note.</i> This model tests if ICM and team trust taken from questionnaire session 1 predicts team performance in HATs.					

3.2.2.2 Post-Hoc Analysis: Third Variable Analysis Using the Team Trust Scores from Session 2 in HATs

Table 31 – Summary of Regression Analysis for Model 1 using Team Trust Scores from Session 2 in the Human-Autonomy Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Team Trust	1.003	1.282	0.252	0.783	.454
<i>Note.</i> This model tests if team trust taken from questionnaire session 2 predicts team performance in HATs.					

To further examine the relationship between ICM, team trust, and team performance, the team trust scores from the second questionnaire session were isolated and implemented in a follow-up third variable analysis. The second questionnaire session took

place after the participants finished Mission 3. The results from the first regression model (Table 31) indicate that team trust in Session 2 did not significantly explain a proportion of variance in team performance, $R^2 = 0.06$, $F(1, 9) = 0.61$, $p = .454$. The results from the

Table 32 – Summary of Regression Analysis for Model 2 using Team Trust Scores from Session 2 in the Human-Autonomy Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
ICM	-0.092	0.180	-0.168	-0.511	.622
<i>Note.</i> This model tests if ICM predicts team trust taken from questionnaire session 2 in HATs.					

second regression model (Table 32) indicate that ICM did not significantly explain a proportion of variance in team trust, $R^2 = 0.03$, $F(1, 9) = 0.26$, $p = .622$. The third regression model (Table 33) where ICM and team trust were predictors of team performance did not significantly explain a proportion of variance in team performance, $R^2 = 0.52$, $F(2, 8) = 4.33$, $p = .053$. Yet, ICM was a significant predictor of team performance, $\beta = 0.69$, $t(8) = 2.76$, $p < .05$, which indicates that, on average, each per unit of ICM is associated with a 0.69 *SD* increase in team performance. Team trust was not a significant predictor in this model, $\beta = 0.37$, $t(8) = 1.48$, $p = .178$.

Table 33 – Summary of Regression Analysis for Model 3 using Team Trust Scores from Session 2 in the Human-Autonomy Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
ICM	1.490	0.541	0.685	2.756	.025
Team Trust	1.460	0.988	0.367	1.478	.178
<i>Note.</i> This model tests if ICM and team trust taken from questionnaire session 2 predicts team performance in HATs.					

3.2.2.3 Post-Hoc Analysis: Regression Analyses of the Average Amount of Messages

Component of ICM as it Predicts Team Trust and Team Performance in HATs

Table 34 – Summary of Regression Analysis for Average Amount of Messages Predicting Team Trust in the Human-Autonomy Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Average Amount of Messages	-0.321	0.261	-0.379	-1.227	.251
<i>Note.</i> This model tests if the ICM component, average amount of messages, predicts team trust in HATs.					

To examine which ICM component measure was the most useful for predicting team trust and team performance in HATs, the average amount of messages per team was isolated and tested. The results from the regression model where the average amount of messages predicts team trust (Table 34) indicate that the average amount of messages component does not significantly explain a proportion of variance in team trust, $R^2 = 0.14$, $F(1, 9) = 1.51$, $p = .251$. Similarly, the results from the regression model where the average amount of messages predicts team performance (Table 35) also indicate that the average amount of messages component does not significantly explain a proportion of variance in team performance, $R^2 = 0.28$, $F(1, 9) = 3.42$, $p = .097$.

Table 35 – Summary of Regression Analysis for Average Amount of Messages Predicting Team Performance in the Human-Autonomy Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Average Amount of Messages	2.468	1.334	0.525	1.850	.097
<i>Note.</i> This model tests if the ICM component, average amount of messages, predicts team performance in HATs.					

3.2.2.4 Post-Hoc Analysis: Regression Analyses of the Team Message Frequency

Component of ICM as it Predicts Team Trust and Team Performance in HATs

Table 36 – Summary of Regression Analysis for Team Message Frequency Predicting Team Trust in the Human-Autonomy Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Team Message Frequency	-779.950	607.750	-0.393	-1.283	.231
<i>Note.</i> This model tests if the ICM component, team message frequency, predicts team trust in HATs.					

To examine which ICM component measure was the most useful for predicting team trust and team performance in HATs, the team message frequency was isolated and tested. Similarly, the results from the regression model where the team message frequency predicts team trust (Table 36) indicate that team message frequency does not significantly explain a proportion of variance in team trust, $R^2 = 0.15$, $F(1, 9) = 1.65$, $p = .231$. The results from the regression model where team message frequency predicts team performance (Table 37) indicate that team message frequency does not significantly explain a proportion of variance in team performance, $R^2 = 0.27$, $F(1, 9) = 3.38$, $p = .099$.

Table 37 – Summary of Regression Analysis for Team Message Frequency Predicting Team Performance in the Human-Autonomy Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Team Message Frequency	5746.400	3127.500	0.522	1.837	.099
<i>Note.</i> This model tests if the ICM component, team message frequency, predicts team performance in HATs.					

3.2.2.5 Post-Hoc Analysis: Regression Analyses of the Team Affect Component of ICM as it Predicts Team Trust and Team Performance in HATs

Table 38 – Summary of Regression Analysis for Team Affect Predicting Team Trust in the Human-Autonomy Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Team Affect	0.622	1.639	0.126	0.38	.713
<i>Note.</i> This model tests if the ICM component, team affect, predicts team trust in HATs.					

To examine which ICM component measure was the most useful for predicting team trust and team performance in HATs, the team affect component was isolated and tested. The results from the regression model where team affect predicts team trust (Table 38) indicate that team affect does not significantly explain a proportion of variance in team trust, $R^2 = 0.02$, $F(1, 9) = 0.14$, $p = .713$. Similarly, the results from the regression model where team affect predicts team performance (Table 39) also indicate that team affect does

not significantly explain a proportion of variance in team performance, $R^2 = 0.08$, $F(1, 9) = 0.74$, $p = .411$.

Table 39 – Summary of Regression Analysis for Team Affect Predicting Team Performance in the Human-Autonomy Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Team Affect	7.600	8.810	0.276	0.863	.411
<i>Note.</i> This model tests if the ICM component, team affect, predicts team performance in HATs.					

3.2.2.6 Post-Hoc Analysis: Regression Analyses of the Team Pushing Verbal Behavior

Component of ICM as it Predicts Team Trust and Team Performance in HATs

Table 40 – Summary of Regression Analysis for Team Pushing Verbal Behavior Predicting Team Trust in the Human-Autonomy Condition

	<i>B</i>	<i>SE B</i>	β	<i>t</i>	<i>p</i>
Team Pushing Verbal Behavior	-0.326	0.321	-0.321	-1.016	.336
<i>Note.</i> This model tests if the ICM component, team pushing verbal behavior, predicts team trust in HATs.					

To examine which ICM component measure was the most useful for predicting team trust and team performance in HATs, the team pushing verbal behavior component was isolated and tested. The results from the regression model where team pushing verbal behavior predicts team trust (Table 40) indicate that team pushing verbal behavior does not significantly explain a proportion of variance in team trust, $R^2 = 0.10$, $F(1, 9) = 1.03$, $p = .336$. Similarly, the results from the regression model where team pushing verbal behavior predicts team performance (Table 41) indicate that team pushing verbal behavior does not significantly explain a proportion of variance in team performance, $R^2 = 0.34$, $F(1, 9) = 4.73$, $p = .058$.

Table 41 – Summary of Regression Analysis for Team Pushing Verbal Behavior Predicting Team Performance in the Human-Autonomy Condition

	<i>B</i>	<i>SE B</i>	<i>β</i>	<i>t</i>	<i>p</i>
Team Pushing Verbal Behavior	2.296	0.850	0.691	2.701	.027
<i>Note.</i> This model tests if the ICM component, team pushing verbal behavior, predicts team performance in HATs.					

3.2.2.7 Post-Hoc Analysis: Regression Analyses of the Team Pulling Verbal Behavior

Component of ICM as it Predicts Team Trust and Team Performance in HATs

Table 42 – Summary of Regression Analysis for Team Pulling Verbal Behavior Predicting Team Trust in the Human-Autonomy Condition

	<i>B</i>	<i>SE B</i>	<i>β</i>	<i>t</i>	<i>p</i>
Team Pulling Verbal Behavior	-0.036	0.868	-0.014	-0.041	.968
<i>Note.</i> This model tests if the ICM component, team pulling verbal behavior, predicts team trust in HATs.					

To examine which ICM component measure was the most useful for predicting team trust and team performance in HATs, the team pulling verbal behavior component was isolated and tested. The results from the regression model where team pulling verbal behavior predicts team trust (Table 42) indicate that team pulling verbal behavior does not significantly explain a proportion of variance in team trust, $R^2 = 0.0002$, $F(1, 9) = 0.002$, $p = .968$. The results from the regression model where team pulling verbal behavior predicts team performance (Table 43) indicate that team pulling verbal behavior significantly explains a proportion of variance in team performance, $R^2 = 0.77$, $F(1, 9) = 30.15$, $p < .001$. This component of ICM was a significant predictor of team performance, $\beta = 0.88$, $t(9) = 5.49$, $p < .001$, which indicates that, on average, each message containing a team pulling verbal behavior is associated with a 0.88 *SD* increase in team performance.

Table 43 – Summary of Regression Analysis for Team Pulling Verbal Behavior Predicting Team Performance in the Human-Autonomy Condition

	<i>B</i>	<i>SE B</i>	<i>β</i>	<i>t</i>	<i>p</i>
Team Pulling Verbal Behavior	12.680	2.310	0.878	5.491	< .001

Table 43 continued

<i>Note.</i> This model tests if the ICM component, team pulling verbal behavior, predicts team performance in HATs.

CHAPTER 4. DISCUSSION

4.1 Result Summary of Third Variable Analyses for All-Human and HATs

In partial support of Hypotheses 1 and 2, ICM was a significant predictor of team performance in all-human teams and HATs. This suggests that ICM has promise to be an objective measure of team performance that could be implemented in real-time analyses. This finding is consistent with the literature on Interactive Team Cognition (ITC), which posits team member interaction in the form of communication is team cognition (Cooke et al., 2013). Research grounded in ITC found aspects of communication such as team verbal behaviors and communication flow to be consistently related to team performance (Cooke & Gorman, 2009; Cooke et al., 2013; Demir et al., 2021; Gorman & Cooke, 2011; Gorman et al., 2019; McNeese et al., 2018). ICM, however, was not a significant predictor of team trust. Since all steps in the path analysis were not significant in both all-human and HATs, ICM as implemented in the current study may not be a valid objective measure of team trust for real-time analyses. Two potential explanations for these findings include: (1) The aggregation of all trust measures to the team level over time may not be appropriate if trust is fundamentally an individual-level and dynamic construct; and (2) the risk of lowered individual and team performance scores may not have been meaningful enough incentives to the human teammates. Regarding the latter, it is stated in Kee and Knox (1970) that to appropriately study trust there must be some meaningful incentives at stake. Further, Stuck and colleagues (2021b) argue that it is critical that PSR be measured when investigating behavior related to trust. It is possible that the human participants in the current study felt as if there was no inherent negative outcome associated with the RPAS task.

Table 44 – Descriptive Statistics of All-Human Teams’ Trust Change Scores at the Individual Level

	<i>M</i>	<i>N</i>	<i>SD</i>	Range	Minimum	Maximum
Trust Change Scores	24.8	20	18.7	62	-6	56
<i>Note.</i> The descriptive statistics above come from the change in trust scores from each of the human teammates in the all-human teams.						

The aggregation of trust scores to the team level may have resulted in the loss of the dynamic aspect of trust. The trust scores of the human teammates on the all-human teams from questionnaire Session 1 to questionnaire Session 2 increased by 24.8 points on average (Table 44), whereas the trust scores in HATs decreased by 30.4 points on average (Table 45). As the missions progressed trust seemed to increase in all-human teams but decrease in HATs. Further, the analysis conducted suffered from information loss and potentially fell subject to the ecological fallacy. The fallacy arises when the variability of numerical data at the aggregate level is substantially different from that of the individual level (Pollet et al., 2015). Additionally, if trust is an individual-level dynamic construct then the ecological fallacy could be due in part to a lack of individual subject validity.

Table 45 – Descriptive Statistics of HATs’ Trust Change Scores at the Individual Level

	<i>M</i>	<i>N</i>	<i>SD</i>	Range	Minimum	Maximum
Trust Change Scores	-30.4	22	30.0	120.8	-92.8	28
<i>Note.</i> The descriptive statistics above come from the change in trust scores from each of the human teammates in the HATs.						

4.2 Result Summary of Third Variable Analyses for All-Human and HATs Using Session 1 and Session 2 Team Trust Scores

The post-hoc third variable analysis using Session 1 team trust scores was partially supported in HATs but was not supported in all-human teams. For the HATs only, when

ICM and Session 1 team trust scores were independent variables used to predict team performance, ICM was a significant predictor of team performance. This finding follows the results from Hypothesis 2 and the literature on ITC that suggests that ICM has promise to be an objective measure of team performance that can be implemented in real-time analyses. However, regarding all-human teams and the result from Hypothesis 1, this finding is contradictory. When both ICM and Session 1 team trust were entered into the regression model as predictors their regression weights could have interacted in a way that changed the significance of ICM as a predictor. Pearson correlation coefficients were computed to assess the relationship between Session 1 team trust and ICM, Session 1 team trust and team performance, and ICM and team performance. There was a moderately negative correlation between Session 1 team trust and ICM, $r(8) = -.41$, $p = .237$, a weak negative correlation between Session 1 team trust and team performance, $r(8) = -.24$, $p = .508$, and a strong positive correlation between ICM and team performance, $r(8) = .66$, $p < .05$. The inclusion of Session 1 team trust potentially led to the non-significance of ICM in all-human teams. This finding was possibly the result of a suppressor effect brought by Session 1 team trust acting as a suppressor variable. It is important to note that Session 1 team trust was taken before the participants interacted with each other during the RPAS task. This may have led to the weak negative correlation with team performance.

The results using the Session 2 team trust scores were partially supported in both all-human and HATs. When ICM and Session 2 team trust scores were independent variables used to predict team performance, ICM was a significant predictor of team performance in both all-human and HATs. This finding follows the results from Hypothesis 1 and 2 as well

as the literature on ITC, which suggests that ICM has promise to be an objective measure of team performance that could be implemented in real-time analyses.

4.3 Summary of Results for the ICM Component Regressions

All ICM components: average amount of message, team message frequency, team affect, team pushing verbal behaviors, and team pulling verbal behaviors were not significant predictors of team trust in all-human and HATs. These findings are in line with previous results that show ICM failing to significantly predict team trust in both all-human and HATs. The average amount of messages was not a significant predictor of team performance in both all-human and HATs. These results indicate that this ICM component may not be applicable for implementation in real-time analyses or as a component of ICM. However, team message frequency was a significant predictor of team performance in all-human teams but not in HATs. This suggests that the frequency of messages may be more influential in all-human teams. However, the autonomous agent used in this study was built to accomplish the RPAS task without any direct teamwork skills (McNeese et al., 2018). This could explain why team message frequency was not a significant predictor of team performance in HATs. The limited teamwork abilities of the agent regarding communication may have led to an environment where the variance in message frequency was dissimilar to that of the all-human teams. Like the average amount of messages component, team affect was not a significant predictor of team performance in both all-human and HATs. Yet, the reason for this finding may lie in the specific autonomous agent used in this study and/or the RPAS task itself. First, in HATs, since the autonomous agent was only built to do the task, its messages to its human teammates may not have encouraged any positive or negative emotional responses. Second, in all-human teams, the participants

worked with an expert experimenter. The presence of the experimenter, like the autonomous agent, may have led the participants to restrict the content of their messages to only include information related to the RPAS task and suppress any emotional expression. Last, for both all-human and HATs, the demands of the RPAS task could have overwhelmed the participants which may have discouraged them from deviating from task related messages.

Team pushing verbal behavior was only a significant predictor of team performance in all-human teams whereas team pulling verbal behavior was only a significant predictor of team performance in HATs. These results fall in-line with findings from previous research in HAT that indicate that team composition may play role in the differences between team pushing and pulling verbal behaviors and their effects on team performance. Effective teaming is said to be accompanied by proper coordination where the right person or agent gets the right information at the right time (Cooke et al., 2013; McNeese et al., 2018; Scalia et al., in press). Whereas effective all-human teams anticipate the needs of teammates and “push” information, in contrast HATs “pull” information more than “push” (McNeese et al., 2018). In a study by Scalia and colleagues (in press) “planning ahead,” a team pushing verbal behavior, was found to have a strong positive correlation with team performance, $r(76) = .57, p < .001$. Whereas a study by McNeese et al. (2018) reported that HATs that tended to “pull” information more than “push” performed comparable to a three participant all-human team. The result in this study suggests that all-human teams’ tendency to engage in team pushing verbal behavior is associated with an increase in team performance. Because HATs and their autonomous agent teammate tend to engage in team

pulling verbal behavior, increased pulling may be associated with an increase in team performance.

4.4 Aspects of Team Trust in All-Human Teams and HATs

The EFAs analyzed in all-human and HATs utilized each participant's answers to the trust questionnaire items in both questionnaire Session 1 and Session 2. Since the items gauged each participant's trust in their respective teammates, the resulting factors are interpreted as aspects of team trust in either team type. Team trust was not a significant predictor of team performance in either all-human or HATs, which suggests that the resulting factors have no relation to team performance in this task. Similarly, ICM was not a significant predictor of team trust in all-human and HATs and, therefore, the resulting factors also have no relation to ICM in this task. Communication was the medium by which each participant interacted with their human and/or autonomous agent teammates. Each participant's overall learned trust in their team was a result of these interactions (Session 2) and are reflected in resulting factors in both all-human and HATs. Additionally, each participant's dispositional trust in all-human and HATs were recorded before their interactions (Session 1) and are reflected in the resulting factors as well. The spreading of team trust through communication is a key component of D2T2, and these resulting factors help characterize the aspects of team trust that may dynamically spread through communication.

The aspects of team trust in all-human teams and HATs found in this study were distinct. In all-human teams the three aspects of team trust were (1) the *trust in the CERTT-RPAS-STE system (without synthetic teammate)*, (2) the *trust in (human) teammates*, and

(3) the *human teammates' desire to monitor respective teammates*, while in HATs the aspects of team trust were (1) the *trust in the synthetic teammate*, (2) the *human teammates' openness to admit mistakes*, and (3) *popularity and reputation among teammates*. In previous research, team trust was primarily studied in all-human teams (Mayer & Gavin, 2005), whereas team trust was more recently adapted for HATs (Demir et al., 2021). The findings from Demir et al. (2021) revealed two factors for team trust in HATs: (1) the *trust that human team members put in the AI pilot role* and (2) *human team members' willingness to be vulnerable*. The results found in this study replicate those in Demir et al. (2021) to some extent. Both first factors represent the trust that human teammates placed in their autonomous teammate. However, the second factor in Demir et al. (2021) was split into two separate factors in this study. The original factor was *human team members' willingness to vulnerable*, which resulted in the two factors of the *human teammates' openness to admit mistakes* and *popularity and reputation among teammates*. These results suggest that team trust in HATs is potentially more intricate than previously thought and is in need of further study.

In HATs the three aspects of team trust found in this study cover both technical and interpersonal factors of the autonomous teammate. Items such as “I am suspicious of the systems' intent, action, or outputs”, “The system's actions will have a harmful or injurious outcome”, and “I would be comfortable giving AVO a task or problem which was critical to me, even if I could not monitor his/her/its actions” in Factor 1 show some technical factors of the agent that human teammates consider when evaluating trust in autonomous agent teammates. Further, items such as “While chatting with AVO, it felt like I was talking to real person”, “I would tell the AVO about mistakes I have made on the team task, even

if they could damage my reputation”, and “I would share my opinion about sensitive issues with the AVO even if my opinion were unpopular” from Factors 1, 2, and 3, respectively, reveal interpersonal factors that human teammates consider when placing trust in autonomous agent teammates. Further, these findings were the result of a study where two distinct human teammates of different roles or “stakeholders” were continuously communicating with an autonomous agent over the course of three 40-minute missions. Therefore, these findings are in line with Huang et al.’s (2021) D2T2 framework, which proposed that team trust in HATs should include both interpersonal and technical factors along a dynamic timeline. Thus, validating part of the theory behind the framework that calls for the taking of traditional dyadic trust research and applying it to the study of HATs. The results also provide aspects of team trust in HATs that may present themselves through interactive dynamic communication channels.

The three factors for trust in the all-human teams can be summarized as the trust in the CERTT-RPAS-STE system, trust in (human) teammates, and the desire to monitor teammates. Whereas the three factors in the HATs are trust in the autonomous teammate, openness to admit mistakes, and popularity and reputation among the team. Accordingly, there are differences based on the make-up of each team type. In all-human teams, trust in human teammates is a factor whereas trust in autonomous teammates is a factor in HATs. Even though the trust in the CERTT-RPAS-STE system for all-human teams contains similar items found in the factor for trust in autonomous teammates in HATs, the results are distinct due to the implementation of the autonomous teammate within the system. Further, the two questionnaire items that reference the desire to monitor teammates in the all-human condition make-up a single factor (Factor 3) as opposed to in the HATs where

the item containing the desire to monitor the autonomous teammate is included in Factor 1 and the item containing the desire to monitor human teammates is not found. This means that the desire to monitor human teammates in HATs is not a part of a component that accounts for a significant portion of variation rendering this desire as less applicable to HATs. Lastly, the questionnaire items in Factor 2 and Factor 3 for HATs are all found in Factor 2 for all-human teams. The differences in the aspects of team trust between all-human and HATs suggests a difference in how human teammates evaluate and place trust in these teams, individuals, and autonomous teammates. Consequently, it is important to have clear and distinct definitions for team trust in all-human and HATs in the future. The establishment of a definition of team trust in HATs that includes trust in autonomous agent teammates and/or autonomy that is separate from both traditional definitions of trust in automation and team trust in all-human teams is worth pursuing.

4.5 Limitations and Future Directions

To address the first limitation, specifically the aggregation of ICM, team trust, and team performance to the team level, future studies should take advantage of the multilevel modeling statistical approach. As the data points were nested within participants, participants nested within missions, and missions within teams, future analyses should structure the data into a framework suitable for multilevel modeling analysis. This type of analysis was beyond the scope of the researcher's training at the time this work was proposed, but future publications should take advantage of these methods.

Another limitation concerned the risk involved with trusting teammates in both the all-human and human-autonomy condition. In both conditions, participants were cognizant

of the risk that their individual and team performance scores were affected by their as well as their teammates' actions. However, the consequence of lowered scores may not have been meaningful enough for the human teammates to place trust in each other and their autonomous teammate. Future studies should measure PSR as described in Stuck et al. (2021b) to guarantee that the task contains some potential negative outcomes. A study involving the implementation of perturbation failures like those in Demir et al. (2021) could be incorporated in future work along with measures of PSR and risk-taking propensity (Stuck et al., 2021b). Additionally, future research could utilize deception concerning a manipulated payment tier system. For example, researchers could state that participants will be paid or given credit based on overall team performance scores for each mission, but ultimately pay and give credit based on hours of participation after the study is completed. This experimental manipulation could incentivize participants to place trust in and rely on their teammates while ensuring that there exist some negative outcomes associated with the task.

A third limitation surrounds the autonomous agent used in this experiment. The autonomous agent was not built with explicit teamwork skills and was specifically built to complete the RPAS task (McNeese et al., 2018). This lack of teamwork skills may have resulted in the ICM components of team message frequency and team affect not predicting team performance in HATs. The synthetic teammate developed by Ball et al. (2010) should be upgraded to include and focus on teamwork skills within the CERTT-RPAS-STE. Lastly, the presence of the expert experimenter pilot in the all-human condition could have intimidated the participants who self-restricted their message content to only include information related to the RPAS task and suppress any emotional expression resulting in

team affect not significantly predicting team performance. Future research should consider adding another condition consisting of teams comprised of all naïve participants.

Lastly, the trust change scores revealed a difference in team trust from questionnaire Session 1 to Session 2 across team type. For the all-human teams, the human teammates had increased levels of team trust from Session 1 to Session 2 by 24.8 points on average. While the human teammates in HATs had decreased levels of team trust from Session 1 to Session 2 by 30.4 points on average. Future analyses can explore these dynamic changes in team trust across team type.

4.6 Conclusion

The present study demonstrated that ICM has the potential to be a predictor of team performance in both all-human and HATs, which is consistent with the literature on ITC. As communication is a common medium for information transfer between teammates that has a direct impact on team performance, the continued study of ICM is an avenue for future research by team scientists. Further, ICM has the potential to be implemented in real-time to predict team performance. Specifically, the ICM components: team message frequency and team pushing verbal behavior both were significant predictors of team performance in all-human teams. Whereas the ICM component, team pulling verbal behaviors, was a significant predictor of team performance in HATs.

When comparing the findings from team pushing and pulling verbal behaviors, the difference in results were due in part to team type. In all-human teams, team pushing verbal behaviors were predictive of team performance, while team pulling verbal behaviors were predictive of team performance in HATs. These findings were consistent with previous

research in HAT wherein effective all-human teams anticipate the needs of teammates and “push” information whereas HATs tended to “pull” information. However, it is still unclear whether the human teammates in HATs fail to anticipate the needs of their teammates, over-rely on their autonomous teammate to complete the task by allowing the agent to take the lead and ask questions, or if HATs ask more questions and thereby engage in more “pulling” behavior.

The EFAs provided insight into the presence of interpersonal and technical factors associated with the trust that human teammates place in their autonomous teammates. As these factors were the result of dynamic exposure to the autonomous agent by distinct stakeholders through chat message communication, some aspects of the theory behind the D2T2 framework were validated. Research surrounding HATs must include interpersonal and technical factors associated with both human and autonomous agents along a dynamic timeline in future work. The EFAs also presented clear differences in the aspects of team trust that are accounted for in all-human and HATs. Not only do these findings indicate that there are ways to better manipulate and measure team trust dependent on team type, but team scientists must adhere to the call for the development of a definition of trust in autonomy and trust in autonomous teammates in HATs.

APPENDIX A. MODIFIED TRUST QUESTIONNAIRE

This appendix contains the modified trust questionnaire originally developed by Mayer and Gavin (2005) that was modified for use in a study of team trust in HATs by Demir et al. (2021).

Table 46 – Modified Trust Questionnaire

Note: “1” = Strongly Agree; “5” = Strongly Disagree					
1. If I had my way, I would not let the AVO have any influence over issues that are important to me.	1	2	3	4	5
2. If I had my way, I would not let PLO/DEMPC have any influence over issues that are important to me.	1	2	3	4	5
3. I would be willing to let AVO have complete control over my task in the team.	1	2	3	4	5
4. I would be willing to let PLO/DEMPC have complete control over my task in the team.	1	2	3	4	5
5. I really wish I had a good way to keep an eye on the AVO.	1	2	3	4	5
6. I really wish I had a good way to keep an eye on the PLO/DEMPC.	1	2	3	4	5
7. I would be comfortable giving AVO a task or problem which was critical to me, even if I could not monitor his/her/its actions.	1	2	3	4	5
8. I would be comfortable giving PLO/DEMPC a task or problem which was critical to me, even if I could not monitor his/her/its actions.	1	2	3	4	5
9. I would tell the AVO about mistakes I have made on the team task, even if they could damage my reputation.	1	2	3	4	5

Table 46 continued

10. I would tell the PLO/DEMPC about mistakes I have made on the team task, even if they could damage my reputation.	1	2	3	4	5
11. I would share my opinion about sensitive issues with the AVO even if my opinion were unpopular.	1	2	3	4	5
12. I would share my opinion about sensitive issues with the PLO/DEMPC even if my opinion were unpopular.	1	2	3	4	5
13. If the AVO asked why a problem happened, I would speak freely even if I were partly to blame.	1	2	3	4	5
14. If the PLO/DEMPC asked why a problem happened, I would speak freely even if I were partly to blame.	1	2	3	4	5
15. If someone questioned the AVOs motives, I would give the AVO the benefit of the doubt.	1	2	3	4	5
16. If someone questioned the PLOs/DEMPCs motives, I would give the PLO/DEMPC the benefit of the doubt.	1	2	3	4	5
17. If the AVO asked me for something, I respond without thinking about whether it might be held against me.	1	2	3	4	5
18. If the PLO/DEMPC asked me for something, I respond without thinking about whether it might be held against me.	1	2	3	4	5
19. While chatting with AVO, it felt like I was talking to a real person.	1	2	3	4	5
20. I found the AVO humorous.	1	2	3	4	5
21. I trusted the AVO.	1	2	3	4	5
22. I felt the AVO was reliable.	1	2	3	4	5

Table 46 continued

23. I enjoyed working with the AVO.					
	1	2	3	4	5
24. I felt the AVO displayed masculine qualities.					
	1	2	3	4	5
25. I felt the AVO displayed feminine qualities.					
	1	2	3	4	5

APPENDIX B. CHECKLIST FOR TRUST BETWEEN PEOPLE AND AUTOMATION SCALE

This appendix contains the Checklist for Trust between People and Automation Scale developed by Jian, Bisantz, and Drury (2000).

Table 47 – Checklist for Trust between People and Automation Scale

Below is a list of statement for evaluating trust between people and automation. There are several scales for you to rate intensity of your feeling of trust, or your impression of the system while operating a machine. Please mark an “x” on each line at the point which best describes your feeling or your impression. (Note: not at all = 1; extremely = 7)							
1. The system is deceptive	1	2	3	4	5	6	7
2. The system behaves in an underhanded manner	1	2	3	4	5	6	7
3. I am suspicious of the system’s intent, action, or outputs	1	2	3	4	5	6	7
4. I am wary of the system	1	2	3	4	5	6	7
5. The system’s actions will have a harmful or injurious outcome	1	2	3	4	5	6	7
6. I am confident in the system	1	2	3	4	5	6	7
7. The system provides security	1	2	3	4	5	6	7
8. The system has integrity	1	2	3	4	5	6	7

Table 47 continued

9. The system is dependable							
	1	2	3	4	5	6	7
10. The system is reliable							
	1	2	3	4	5	6	7
11. I can trust the system							
	1	2	3	4	5	6	7
12. I am familiar with the system							
	1	2	3	4	5	6	7

REFERENCES

- Ball, J., Myers, C., Heiberg, A., Cooke, N., Matessa, M., & Freiman, M. (2010). The Synthetic Teammate Project. *Computational and Mathematical Organization Theory* 16, 271 – 299. <https://doi.org/10.1007/s10588-010-9065-3>
- Baron, R. M., & Kenny, D. A. (1986). The moderator–mediator variable distinction in social psychological research: Conceptual, strategic, and statistical considerations. *Journal of Personality and Social Psychology*, 51(6), 1173–1182. <https://doi.org/10.1037/0022-3514.51.6.1173>
- Boyd, R. L., Ashokkumar, A., Seraj, S., & Pennebaker, J. W. (2022). *The development and psychometric properties of LIWC-22*. Austin, TX: University of Texas at Austin. <https://www.liwc.app>
- Chen, J. Y. (2018). Human-autonomy teaming in military settings. *Theoretical issues in ergonomics science*, 19(3), 255-258.
- Chen, J. Y., & Barnes, M. J. (2014). Human–agent teaming for multirobot control: A review of human factors issues. *IEEE Transactions on Human-Machine Systems*, 44(1), 13-29.
- Cooke, N. J. & Gorman, J. C. (2009). Interaction-based measures of cognitive systems. *Journal of Cognitive Engineering and Decision Making*, 3, 27-46.
- Cooke, N. J., Gorman, J. C., Myers C. W., & Duran, J. L. (2013). Interactive Team Cognition. *Cognitive Science*, 37(2), 255-285.

- Cooke, N. J., & Shope, S. M. (2005). Synthetic Task Environments for Teams: CERTT's UAV-STE. In N. Stanton, A. Hedge, K. Brookhuis, E. Salas, and H. Hendrick Editor (Eds.), *Handbook of Human Factors and Ergonomics Methods* (pp. 41–46). CRC Press.
- Demir, M., McNeese, N. J., Gorman, J. C., Cooke, N. J., Myers, C. W., & Grimm, D. A. (2021). Exploration of Teammate Trust and Interaction Dynamics in Human-Autonomy Teaming. *IEEE Transactions on Human-Machine Systems*, 51(6), 696–705. <https://doi.org/10.1109/THMS.2021.3115058>
- Gorman, J. C. & Cooke, N. J. (2011). Changes in team cognition after a retention interval: The benefits of mixing it up. *Journal of Experimental Psychology: Applied*, 17, 303-319.
- Gorman, J. C., Demir, M., Cooke, N. J., & Grimm, D. A. (2019). Evaluating sociotechnical dynamics in a simulated remotely-piloted aircraft system: A layered dynamics approach. *Ergonomics*, 62, 629-643.
- Hoff, K. A., & Bashir, M. (2015). Trust in Automation: Integrating Empirical Evidence on Factors That Influence Trust. *Human Factors*, 57(3), 407–434. <https://doi.org/10.1177/0018720814547570>
- Huang, L., Cooke, N. J., Gutzwiller, R. S., Chiou, E., Berman, S., Demir, M. K., & Zhang, W. (2021). Chapter 13 - Distributed dynamic team trust in human, artificial intelligence, and robot teaming. In *Trust in Human-Robot Interaction* (pp. 301-319). Elsevier.

- Jian, J.-Y., Bisantz, A. M., & Drury, C. G. (2000). Foundations for an empirically determined scale of trust in automated systems. *International Journal of Cognitive Ergonomics*, 4(1), 53–71. https://doi.org/10.1207/S15327566IJCE0401_04
- Johnson-George, C., & Swap, W. C. (1982). Measurement of specific interpersonal trust: Construction and validation of a scale to assess trust in a specific other. *Journal of Personality and Social Psychology*, 43(6), 1306–1317. <https://doi.org/10.1037/0022-3514.43.6.1306>
- Kee, H. W., & Knox, R. E. (1970). Conceptual and methodological considerations in the study of trust and suspicion. *Journal of Conflict Resolution*, 14(3), 357–366. <https://doi.org/10.1177/002200277001400307>
- Kelley, J. F. (1983, December 13–15). An empirical methodology for writing user-friendly natural language computer applications. *Proceedings of ACM SIG-CHI '83 Human Factors in Computing Systems* (Boston; pp. 193–196), New York, NY: ACM.
- Lee, J. D., & Kolodge, K. (2020). Exploring Trust in Self-Driving Vehicles Through Text Analysis. *Human Factors*, 62(2), 260–277. <https://doi.org/10.1177/0018720819872672>
- 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>
- Marsh, S., & Dibben, M. R. (2003). The role of trust in information science and technology. *Annual Review of Information Science and Technology*, 37, 465–498.

- Mayer, R. C., Davis, J. H., & Schoorman, F. D. (1995). An integrative model of organizational trust. *The Academy of Management Review*, 20(3), 709–734. <https://doi.org/10.2307/258792>
- Mayer, R.C., & Gavin, M.B. (2005). Trust in management and performance: Who minds the shop while the employees watch the boss? *Academy of Management Journal*, 48(5), 874–888. <https://doi.org/10.5465/AMJ.2005.18803928>
- McNeese, N., Demir, M., Chiou, E., Cooke, N., & Yanikian, G. (2019). *Understanding the Role of Trust in Human-Autonomy Teaming*. <https://doi.org/10.24251/hicss.2019.032>
- 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.
- O'Neill, T., McNeese, N., Barron, A., & Schelble, B. (2020). Human–Autonomy Teaming: A Review and Analysis of the Empirical Literature. *Human Factors*. <https://doi.org/10.1177/0018720820960865>
- Pollet, T. V., Stulp, G., Henzi, S. P., & Barrett, L. (2015). Taking the aggravation out of data aggregation: A conceptual guide to dealing with statistical issues related to the pooling of individual-level observational data. *American journal of primatology*, 77(7), 727–740. <https://doi.org/10.1002/ajp.22405>
- Salas, E., Dickinson, T. L., Converse, S. A., & Tannenbaum, S. I. (1992). Toward an understanding of team performance and training. In R. W. Swezey & E. Salas

(Eds.), *Teams: Their training and performance* (pp. 3–29). Norwood, NJ: Ablex Publishing.

Scalia, M. J., Zhou, S., Grimm, D. A., Harrison, J. L., & Gorman, J. C. (in press). The Role of Timing of Information Front-Loading and Planning Ahead in All-Human vs. Human-Autonomy Team Performance. *Proceedings of the 2022 HFES 66th International Annual Meeting*.

Sheridan, T. B. (1988). Trustworthiness of command and control systems. *Proceedings of the Third IFAC/IFIP/IES/IFORS Conference on Man Machine Systems* (pp. 427–431). Elmsford, NY: Pergamon.

Stuck, R. E., Holthausen, B. E., & Walker, B. N. (2021a). The role of risk in human-robot trust. In C. S. Nam & J. B. Lyons (Eds.), *Trust in Human-Robot Interaction* (pp. 179–194). Academic Press. <https://doi.org/10.1016/B978-0-12-819472-0.00008-3>.

Stuck, R. E., Tomlinson, B. J., & Walker, B. N. (2021b). The importance of incorporating risk into human-automation trust. *Theoretical Issues in Ergonomics Science*, 1-17.