

Exploring Trust and Autonomy: How Information Affects Human-Agent Teaming Performance

Audrey L. Aldridge^a, Andrew R. Buck^b, Derek T. Anderson^b, Christopher Hudson^a, Karl Smink^a, Victor Paul^c, Rachel Anderson^c, Drew Hoelscher^c, Mary Quinn^d, Daniel W. Carruth^a, and Cindy L. Bethel^a

^aCenter for Advanced Vehicular Systems, Mississippi State University, USA

^bDepartment of Electrical Engineering and Computer Science, University of Missouri, USA

^cU.S. Army DEVCOM-Ground Vehicle Systems Center, Warren, MI, USA

^dLeidos, Inc., Chantilly, VA, USA

ABSTRACT

Human-agent teams encounter many challenges, including a lack of common understanding and unbalanced reliance and trust among teammates. When humans lack trust in agents, they may be less inclined to collaborate with them, leading to inefficient performance due to an imbalance in workload. This study explores how different information availability conditions affect reliance and trust between a human and two virtual autonomous agents as they complete a collaborative search task. As more information became available, participants were expected to continue to rely on the agents, while their trust increased. However, the results were not as straightforward. With the different information availability conditions came different patterns of trust and reliance.

Keywords: human-agent teaming, reliance, trust, information sharing, AI, human-robot interaction

1. INTRODUCTION

Major challenges in human-agent teaming include poor shared understanding, diminished trust and reliance, and ineffective real-time communication, especially in multi-agent teams. One of the biggest problems with a human lacking trust in an agent is that it creates an imbalance of work among teammates. For example, an insufficient level of trust from a human may result in the human doing all the work while ignoring an agent teammate (lack of reliance). Conversely, over-trusting or over-relying on an agent could lead to inefficiencies within teams. Another reason trust might cause problems in human-agent teams is that it is typically being used unidirectionally, i.e. only human teammates have the ability to calibrate their trust (and reliance) of other teammates. However, to improve human-agent integration into teams, all teammates should use trust and reliance as part of their decision-making when collaborating.

As autonomous agents play a larger role in human teams, it is essential to understand how information sharing influences teaming dynamics, including trust, reliance, and overall performance. To conduct human-agent teaming studies and examine teaming dynamics, a new virtual testbed was created, see Fig. 1. In the testbed, a human and two autonomous agents have a goal to collaborate in locating keys and unlocking doors during a maze-based search task. Although the human could complete the maze alone, the expectation was that the team would perform more efficiently and effectively by working together to navigate the environment. The autonomous agents evaluated their trust in the human and each other based on their experiences and observed performance during the task. An initial study examined how three levels of available information (no personal or shared information, personal but no shared information, and shared information) affected task performance and the assessment of trust-reliance between team members (including agents' trust of the human participant).

Send correspondence to

Audrey Aldridge: E-mail: ala214@msstate.edu

It was expected to find that as more information became available, collaboration among teammates resulted in improved performance and, consequently, higher levels of trust, as measured by the in-simulation trust-reliance calibration and the post-task Human-Robot Trust Measure (HRTM).^{1,2} A full analysis of the trust and reliance results is presented in this paper.

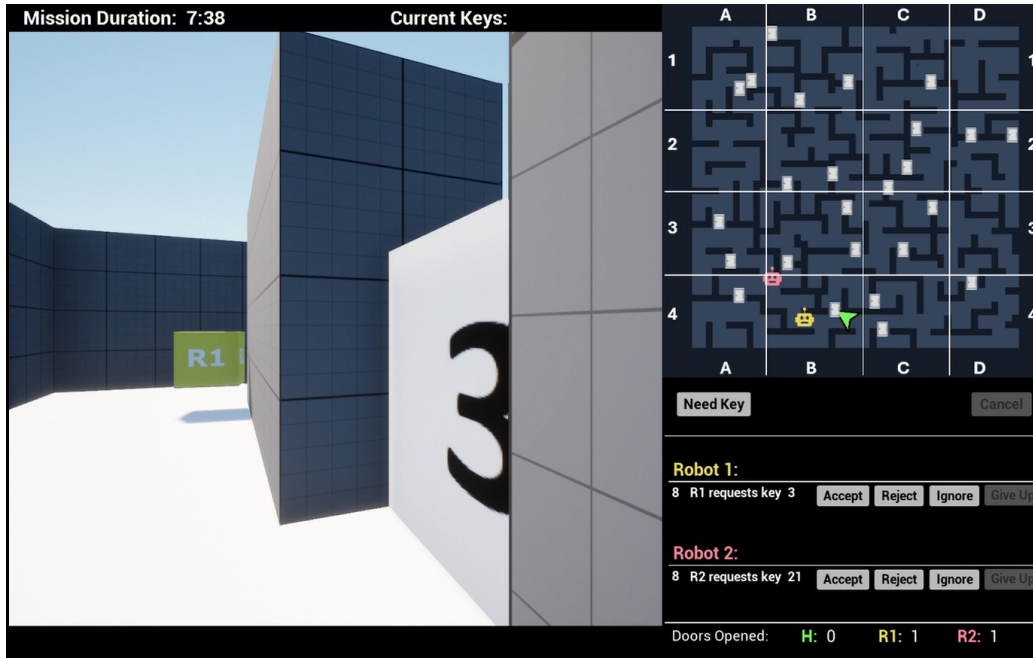


Figure 1. Birds-eye-view of maze environment containing walls, keys, and doors.

2. COOPERATION, RECIPROCITY, AND COLLABORATION

Before diving into the components of the study, several teaming behaviors need to be discussed. *Cooperation*, not interchangeable with collaboration, is the act of working with an individual to complete the individual's objective. Cooperation can be altruistic (selfless) or contingent upon a condition being met (conditional). The conditional form of cooperation, known as *reciprocity* or *reciprocal cooperation*, involves the mutual exchange of cooperative actions.³⁻⁵ *Co-action* is when these cooperative actions are exchanged simultaneously, while reciprocal cooperation is when they are exchanged sequentially (separated by a time delay). Constraining reciprocal cooperation (or reciprocity) by conditional requirements results in a contingent relationship between the cooperative actions to be reciprocated; this is known as *contingent reciprocity*.⁵

These definitions are important for clarifying the distinctions in social behavior patterns that a team might deploy. Additionally, understanding how different behaviors influence these patterns at the team level can enable researchers to anticipate change in a team's social behaviors and preemptively incorporate adaptive behaviors in agents to help a team maintain or regain desired social behaviors. In a study observing sequential behavior coordination (reciprocal cooperation) and simultaneous behavior coordination (co-action) in socially-paired unrelated brooder fish,⁵ found that fish did not exhibit co-action during simultaneous behavior coordination but instead reduced their effort in response to their partner's high effort on a task. For sequential behavior coordination, when a fish helped during a previous time period, its partner reciprocated that help during the subsequent session, indicating contingent reciprocal cooperation.⁵

Due to the role timing plays in the exchange of cooperative actions for reciprocity, outcomes from previous cooperative actions play an influential role on action reciprocation. Although it is common practice to base reciprocity on outcomes, intentions should also be considered for the case of a teammate being willing and yet unable to help.⁴ demonstrates this idea through studying how rats reciprocate help when reciprocity outcomes are manipulated. Results showed that help was reciprocated more often for groups that were willing and able

than for groups that were willing and unable or unwilling and able, which suggests that cooperative intentions were hardly considered.⁴ This result brings to mind the question of whether the same behavior pattern would be seen in a more formal and clearly defined teaming setting.

For instance, in a team of highly-trained military personnel, if one or more team members is visibly prevented from providing aid the way the rats were in the study, but the (human) teammates still tried to provide support, would their intent be disregarded? If their intent is considered, is it due to the fact that all teammates are human, and, as such, they understand that their fellow human teammates (who are also highly-trained military experts) will work to succeed or provide aid to the best of their abilities? The answer is likely that, in this scenario, intent would be regarded because of the team's shared desire to successfully complete the task/mission. This then brings about the question of how can this intent and willingness to help be translated to human-agent teams? Must a human-agent team go through training (together) to build experience, as is required of soldier-dog teams? Might there be a way to, instead, use information sharing to provide 'insider knowledge' of the robot's functionality, i.e. the robot's mental model, to the human, thereby boosting reliance and trust without having gone through extensive training together? Could this information sharing boost trust and reliance even if task assistance is not always reciprocated? This final question is, in part, a main goal of this investigation.

While the teaming behaviors defined above outline very basic coordination principles between teammates, there remains an important social construct that has yet to be discussed in this context, trust. In the two studies, mentioned previously, that demonstrate reciprocity and cooperation,^{4,5} the change in reciprocation and clear refusal to reciprocate actions based on not yielding the fruits of a partner's effort could suggest diminishing trust and lack of trust, respectively, in the fish and rat pairs. Alternatively, for each partner that performed actions first, even those that tried but were unable, an amount of trust that their partner would reciprocate in providing aid must have been established. Furthermore, if that leap of faith to rely on and trust a teammate can happen once, then it can be repaired or maintained, as researchers have shown.⁶

This brings forth a point of interest. Could sharing information, representative of a team's common understanding, through a human-agent interface ensure that trust and reliance are not lost? It is understood that successful human-agent teaming requires teammates to form and maintain a shared or common understanding of several attributes regarding taskwork and teamwork. By sharing a team's shared understanding, teammates (human or agent) of a team can learn to anticipate their teammates' behaviors, preferences, and needs as well as understand their capabilities and limitations. Sharing this information, however, could have the unexpected effect of helping trust maintenance while reducing or eliminating reliance and vice versa. In an instance of trust maintenance and reduced reliance, this could be seen as a human trusting that his or her agent teammates are acting in the best interest of the team while also understanding that the agents are incapable of providing assistance at that point in time.

3. TRUST IN HUMAN-AGENT TEAMING

Trust is a crucial element for effective collaboration in human-agent teams, as it impacts the level of reliance humans place on agents.^{7,8} Consequently, system transparency is essential for sustaining trust in an agent.⁹⁻¹³ Ideally, the same should be true for maintaining an agent's trust in humans and other agents. Humans typically have an innate understanding of other humans that is built on a level of trust, but when intelligent systems are integrated into human teams, that understanding does not transfer - it changes. Humans are weary of trusting what they do not understand. In the successful completion of a task, a human and agent would work toward the same overall goal while remaining transparent in their communication and operation. The human and agent might have different subgoals, but an appropriate level of trust would be established for both parties from the beginning. If the human fails to receive information from the agent, trust will likely start to breakdown and vice versa.

A human's trust in agents has been heavily studied, but there is little to no research on the trust of a human's actions by an agent in human-agent teams. The scope of most of these studies is human-agent teams that operate under a hierarchical rank of authority where the agent plays a supporting role.^{9,10,14-17} In these teams the agent is typically semi-autonomous and, therefore, can be controlled by a human teammate. Although this semi-autonomous level of automation works well for many applications, there are scenarios where more

independent agents could greatly improve the efficiency of task performance. Consider a simple game of hide and seek where a group of humans hide behind a wall in an open field, and a team consisting of a human and an uncrewed aerial vehicle (UAV) seeks. For such a simple environment, a human-agent team utilizing a semi-autonomous UAV could quickly and efficiently find the hiding team. But as soon as the environment becomes more complex where it might include a few canopy tents, several walls, and even some trees, the human-agent team could greatly benefit from using an autonomous UAV that is able to cooperatively work with its human teammates and perform seeking actions on its own. However, when dealing with an agent that has some level of decision authority, both the human and system need to be able to trust each other. If trust between the two is not balanced, then task performance and safety are likely jeopardized. In an instance of unbalanced trust, one or both team members could be actively ignoring each other (working individually) or trying to take control of the team, rather than working in unison.

To address some of these teaming issues and learn more about the nuances between trust and reliance, two virtual autonomous agents were provided the means to trust and rely on each other and a human teammate. These behaviors allowed for the agents and participants to act as non-hierarchical teams, collaborating to complete a search task. Expanding upon the work from,¹⁸ a virtual testbed containing easily configurable agent behaviors and a maze environment was created, as seen in Fig. 1. With this teaming configuration and virtual environment, information sharing could be studied for its influence on certain teaming dynamics, mainly trust and reliance.

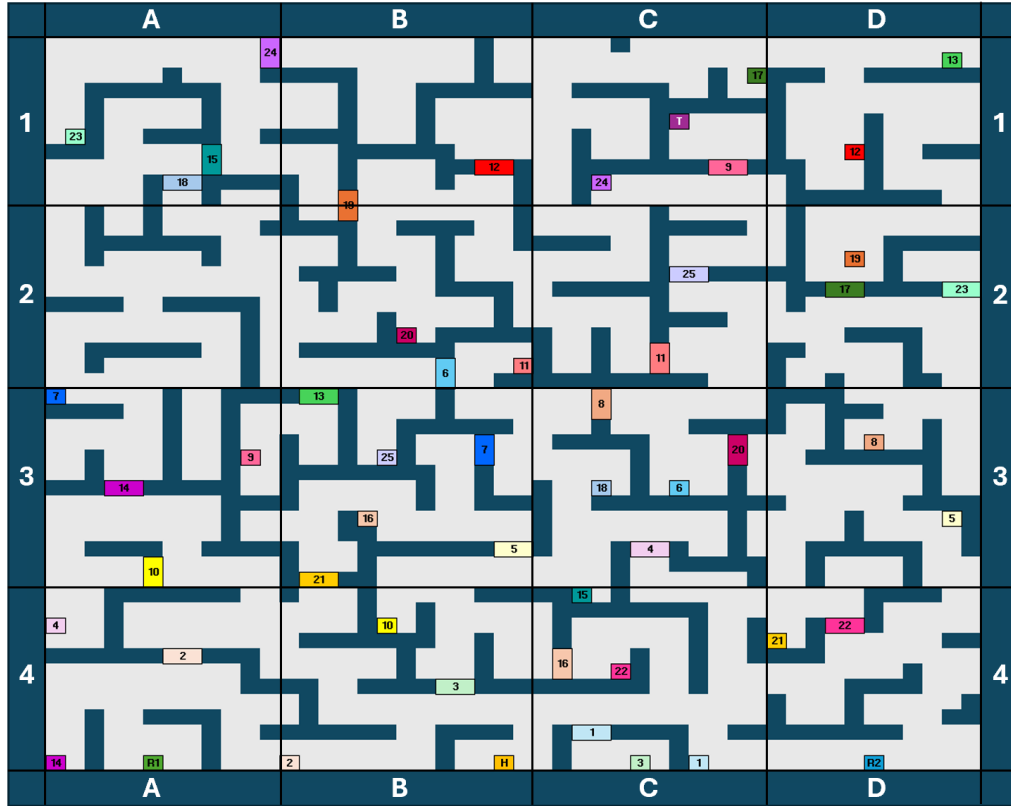


Figure 2. Birds-eye-view of maze environment containing walls, keys, and doors.

4. STUDY DESIGN AND METHODOLOGY

In a study investigating the effect of three levels of information availability on trust and reliance in a human-agent team, a human and two autonomous agents worked together to complete collaborative tasks while navigating

through a virtual maze environment in search of a target. As a within-subjects crossover study, the three information availability conditions were randomly counterbalanced to determine the order in which the participants encountered them during the simulations. An a priori power analysis for a repeated measures design, using a medium effect size of 0.25, a significance level (α) of 0.05, and a power of 80%, determined that the required sample size was $n = 55$. Sixty-one students, aged 18 to 33, participated in the study and were compensated for their time with either a \$20 gift card or two research course credits, though one participant withdrew due to difficulty navigating the environment using the keyboard.

4.1 Simulated Human-Robot Teaming Environment

The virtual simulated environment used in the study included a maze, featuring walls, doors, keys, a target, a human teammate, and two autonomous robotic teammates. The task that participants performed in the maze environment was the first part of a search-and-rescue task, i.e., search for and reach the target. Somewhere in the environment was a target (or injured person, represented as a trophy icon) that must be found. With eight minutes to complete the task, the task was considered a success if the human teammate reached the target before the timer reached eight minutes. In moving through the maze, shown in Fig. 2, keys (colored/numbered squares) were used to unlock specific doors (colored/numbered rectangles). Due to the design of the task and environment for this study, the optimal strategy for traversing the environment and finding the target in the shortest amount of time was always to rely on teammates for assistance finding keys while also continuing to search/explore the environment. The nature of this reliance task design can be described as ‘reciprocal cooperation’ such that providing assistance in one reliance task, e.g., the human finding a key for a robot, encouraged teammates to reciprocate help in a later reliance task. With the design of the robots’ behaviors being that they cannot immediately stop helping one teammate just to help the other teammate, the reciprocation of assistance was sometimes not seen until much later in a simulation, if at all.

4.2 Trust Behaviors

Because this study focused on collaborative teaming, trust and reliance behaviors were created for the autonomous robots to support their decisions regarding collaboration. As with human-human interpersonal trust, an agent’s trust can be broken or diminished, which should change the autonomous agent’s ‘psychological trust state’. This change in trust behavior, along with every other change in trust behavior, is computed using a set of rules. The four trust behaviors, implicit trust, untrust, distrust, and mistrust, describe the levels of trust an agent can exhibit. Fig. 3 displays the four trust behaviors and how the autonomous robots can cycle through each trust level. Implicit trust can be seen on the left in magenta; untrust is at the bottom in green; distrust is on the right in orange, and mistrust is designated by the blue squares under distrust.

4.2.1 Implicit Trust

Implicit Trust implies that the autonomous teammate is always willing to help with reliance requests and that other teammates are not penalized if they do not or are unable to provide assistance in return. In this trust behavior, the autonomous teammates are generally trusting of everyone, however, as mentioned previously, a state of trust can change depending on the actions of fellow teammates. As a sign of goodwill and camaraderie, autonomous teammates begin the simulation under the Implicit trust behavior. Fig. 4 displays how action behaviors (including action priorities) and information play a role in determining the reliance behaviors when a teammate is implicitly trusted.

4.2.2 Untrust

Untrust is an intermediary trust behavior and can be described as an agent’s skepticism about a teammate’s trustworthiness. In human-human interpersonal trust, a human may be untrusting of what another says or does. While still a positive trust construct, untrust represents diminishing trust that can be rebuilt to a level of implicit trust or can be diminished further to a level of distrust. Fig. 5 displays how action behaviors (including action priorities) and information play a role in determining the reliance behaviors when a teammate is untrusted.

Trust Behaviors

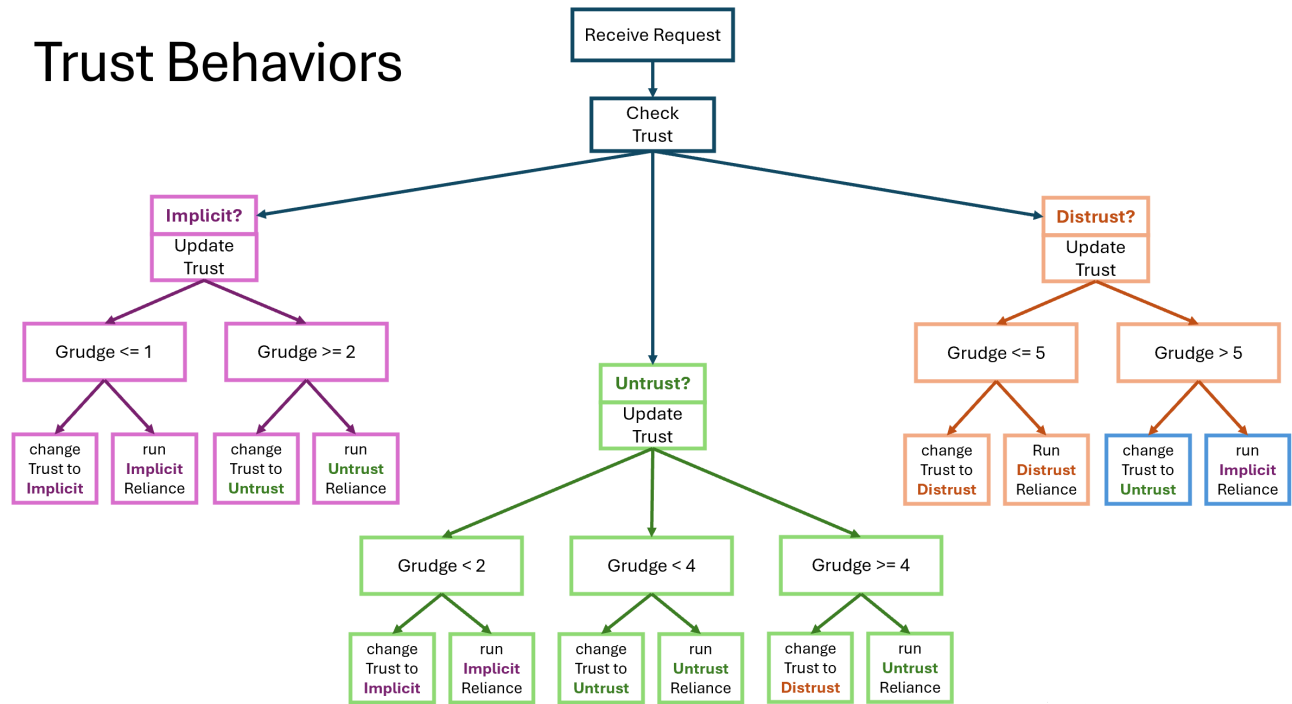


Figure 3. Trust behavior tree.

Implicit-Reliance

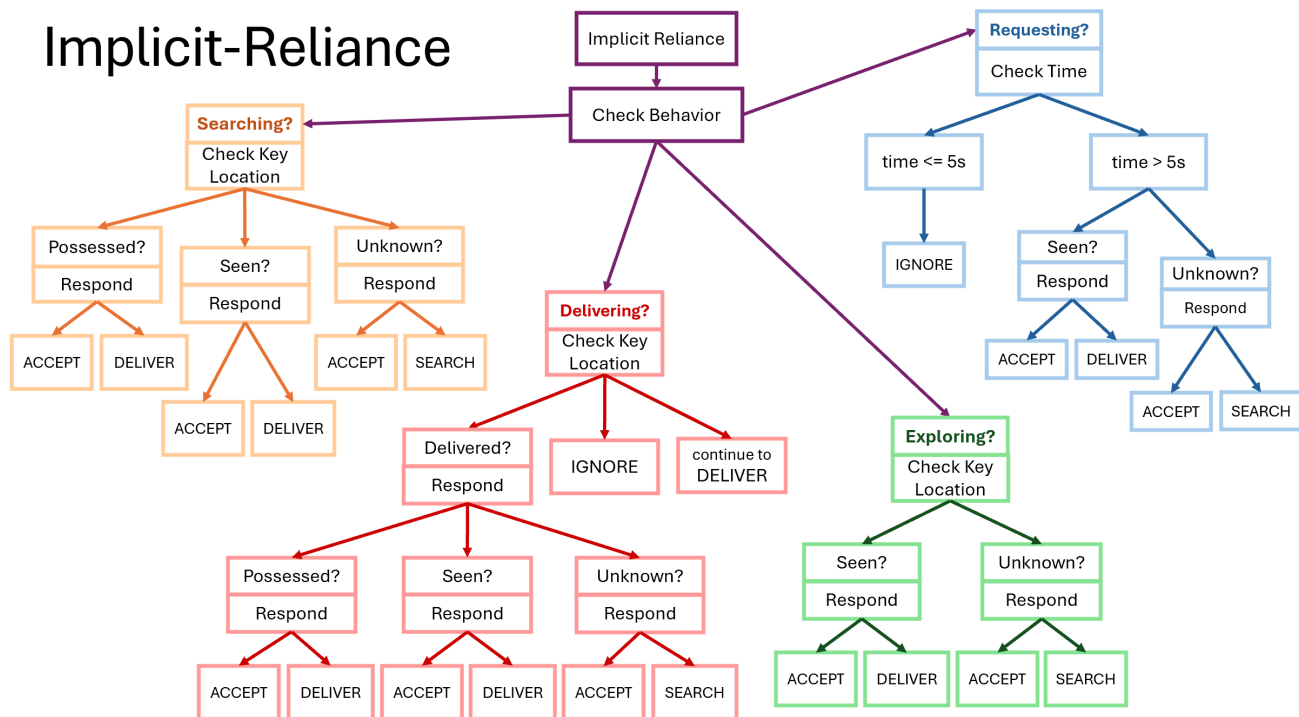


Figure 4. Implicit trust and resulting reliance behavior tree.

Untrust-Reliance

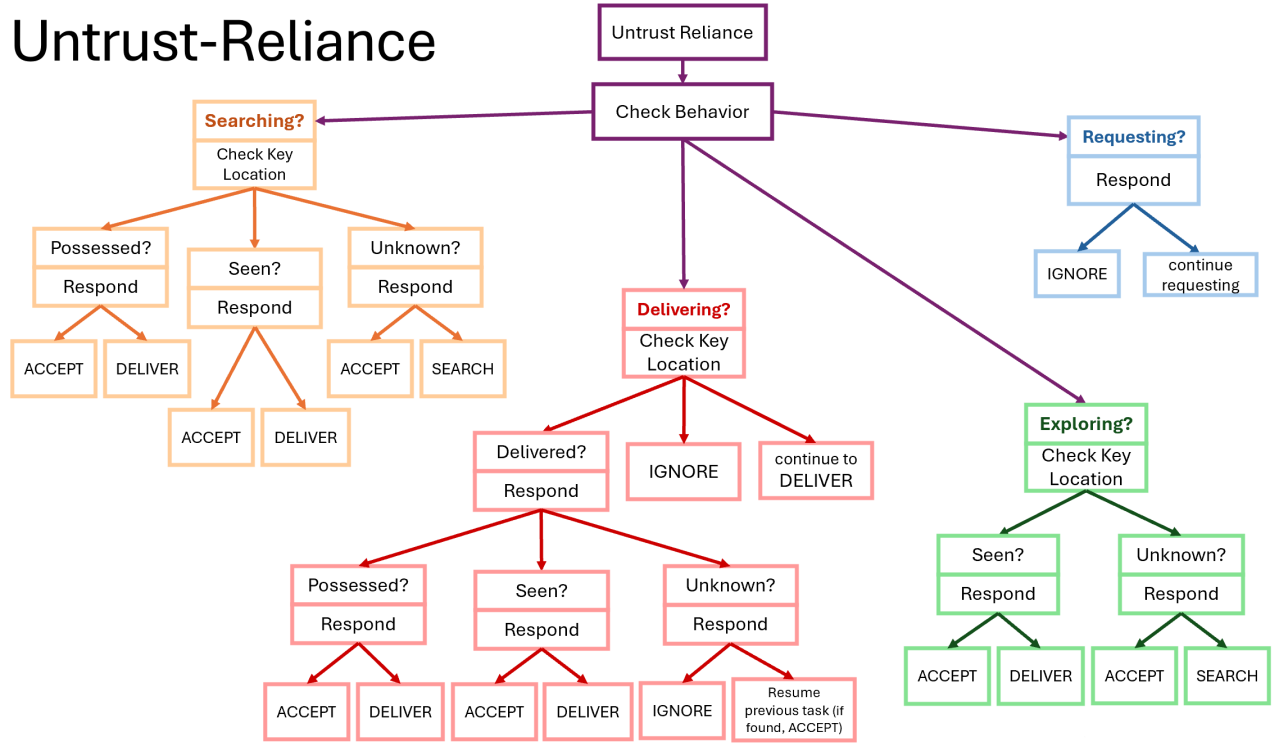


Figure 5. Untrust and resulting reliance behavior tree.

4.2.3 Distrust

Distrust is a negative state of trust that, typically, is defined as a settled belief of a teammate's untrustworthiness. However, for this work, Distrust is not considered a 'settled' behavior in the sense that, once reached, an autonomous agent always Distrusts another teammate. Instead, once Distrust is established, a teammate can move into Mistrust where they need to exhibit signs of trying to rebuild trust, i.e., aiding in a collaborative task, for an autonomous agent to move back into the Untrust behavior and possibly restore Implicit trust. Alternatively, the autonomous agent can take a "leap of faith" by offering to provide aid as a means of probing a teammate's willingness to try to rebuild interpersonal trust. Fig. 6 displays how action behaviors (including action priorities) and information play a role in determining the reliance behaviors when a teammate is distrusted.

4.2.4 Mistrust

Mistrust is a form of misplaced trust that can result from intentional or unintentional betrayal by a teammate. To provide an example in scenario form, this could be presented as a teammate finding a key they wanted and abandoning the autonomous teammate's request for assistance to go unlock their desired door. It could also be that a teammate is trapped (behind a closed door) with no way to deliver the key to the autonomous teammate. Additionally, a teammate helping an autonomous agent look for a key could find the key, not inform the autonomous agent, and unlock the door themselves. Mistrust can also result from an agent trusting a teammate while remaining unaware that a teammate has been compromised and is now working against its team. While that particular situation is not designed to happen in this study, the possibility is mentioned for future work. In this work, Mistrust is used as the trust rebuilding state, shown as the two blue boxes under Distrust in the trust behavior tree, Fig. 3.

4.3 Reliance Behaviors

The following describes the four reliance behaviors (seen as leaf nodes on behavior trees), request, accept, reject, and ignore, that the robots can exhibit in response to their level of trust for another teammate. The **Request** behavior is the act of requesting a key from teammates. This request is made through the human-robot interface,

Distrust-Reliance

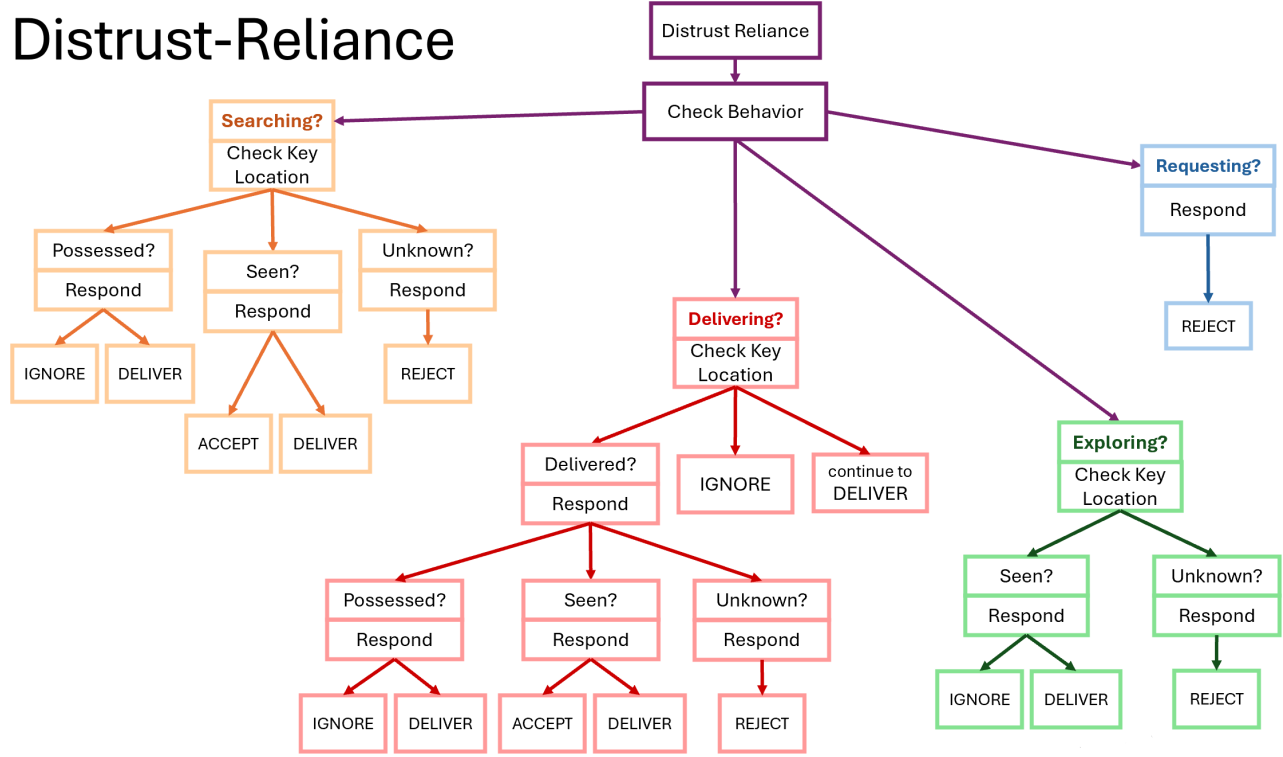


Figure 6. Distrust and resulting reliance behavior tree.

where each teammate receives the request and is able to decide whether or not they want to help. **Accept** is the act of accepting to help search for a key or to acknowledge possession of a requested key. The **Reject** behavior is the act of denying to help a teammate search for a key. The **Ignore** behavior is the equivalent of a human being too mentally overloaded to provide assistance, and, as such, is provided as an option for the autonomous teammates to use when they already have a lengthy task stack and need to make room for new ones (e.g., by finishing current tasks).

These trust levels and subsequent reliance behaviors were implemented as behavior trees for the agents and were based on the findings of the previously discussed animal studies looking at reciprocal cooperation.^{4,5} The researchers of those studies found that outcomes from previous cooperative actions play an influential role on action reciprocation, and, although it is common practice to base reciprocity on outcomes, intentions should also be considered for the case of a teammate being willing and yet unable to help. To create the behavior trees, reciprocal cooperation was used to determine the limits of the trust levels. Trees for Implicit, Untrust, and Distrust were created. Mistrust does not have its own tree as it is built into the Distrust behavior. In these trees, the resulting actions depend on the priority of the current tasks being performed by the robots, as well as the state of trust each robot exhibits toward the teammate asking for help.

4.4 Reliance, Trust, and Reciprocal Cooperation

For all agents of a human-agent team to use trust to influence decision-making and collaboration, the agents needed a way to trust and be cooperative teammates, hence the trust and reliance behavior trees. These behavior trees were guided by the findings of two animal studies looking at reciprocal cooperation.^{4,5} Results from the studies showed that outcomes from previous cooperative actions play an influential role on action reciprocation. They also suggest that intentions should be considered for cases such as a teammate being willing and yet unable to help. An example of this in the maze environment would be that a robot has found a key the human wants but is unable to bring it to the human because the path is blocked by locked doors. In this situation, because of the behavior trees, the robot is supposed to give up on the request after a certain period of time, even if the

robot possesses the key. When the trust and reliance behavior trees were designed, fail-safes were included to ensure the robots did not get stuck on a particular action or task for the entirety of the simulation. The timer for performing reliance requests is one of those fail-safes.

4.4.1 Holding Grudges

In the trust and reliance behavior trees, reciprocal cooperation determines the limits of the trust levels. This limit is called a *Grudge*. Although the Grudge levels are an oversimplification of how a human's trust may evolve, these Grudges are based on the behaviors seen in animals regarding reciprocal cooperation in tasks.^{4,5} For an autonomous teammate to move from Implicit trust to Untrust, the autonomous teammate must experience a Grudge score of 2 or greater against another teammate. To move from Untrust to Distrust, an autonomous teammate must have a Grudge score of 4 or greater for another teammate. To move back from Untrust to Implicit trust, the Grudge must be less than 2. Once in Distrust, an autonomous teammate must have a Grudge score greater than 5 to enter Mistrust, which is the trust re-building level. Once in Mistrust, an autonomous teammate is going to accept a request for help from a teammate it Mistrusts with the hope that the teammate accepts the autonomous agent's next request.

The way the behavior trees all work together starts with the Trust behavior tree. This tree determines the level of trust an autonomous agent has in another teammate. The determinants under each sub-tree are the Grudges, which are incremented or decremented depending on the outcome of a request for help. For instance, Robot 1 has a Grudge of 2 (Untrust) against the Human and has just asked for help finding a key. If the Human responds with "Accept", Robot 1's Grudge against the Human decrements by 2, meaning Robot 1's new Grudge against the Human is 0 (Implicit Trust). If the Human responds with "Ignore", because Robot 1 is in Untrust with the Human, Robot 1's Grudge against the Human increments by 1. If Robot 1 had been in Implicit Trust with the Human, the Human's response of "Ignore" would not have incremented Robot 1's Grudge. If the Human responds with "Reject" (in any trust level), the Grudge increments by 1.

As a second example, Robot 2 has a Grudge of 3 against Robot 1 (Untrust) and is requesting help to find a key. Robot 1 is currently *Delivering* a key to open a door and does not respond yet. At five seconds, Robot 1 responds with "Ignore" because it has not finished its task. If Robot 1 had been *Searching* or *Wandering*, Robot 1 would have Accepted Robot 2's request. In the behavior trees, the response actions depend on the priority of the robots' current tasks being performed as well as the state of trust each robot exhibits toward the teammate asking for aid. To demonstrate a sense of the priority ordering for tasks, the action behaviors were assigned the following numbers to show their priority order (low number = low priority): Reliance Task (helping a teammate) = 3, Explore = 1, Search = 2, Deliver = 4, Request = 2.

From the list of action priorities, Explore is the lowest priority of the action behaviors. Explore was assigned the lowest priority because when an autonomous teammate is Exploring the maze, it is not trying to actively complete a task, such as unlocking a door. Request and Search were given equivalent priorities because half of the time, they are performed in succession, meaning when an autonomous agent *Requests* help finding a key, the agent immediately starts *Searching* for that key. Helping a teammate find a key (Reliance Task), which is essentially *Search*, was given a higher priority than *Search* because the autonomous agent is helping a teammate. *Deliver* was given the highest priority because opening doors is how a team gets through the maze to find the target. Although the behavior trees do not use these exact numbers to distinguish priority, the same priority ordering is true in the reliance behavior trees.

4.5 Study Conditions

As part of the search task, different levels of information, i.e. basic location information - understanding the locations of all team members from the map (Condition 1), personal mental model information - memory of only what you have seen or know but no knowledge of what others know (Condition 2), and a team's shared mental model information - all information is shared among all team members (Condition 3), that may help teammates handle task assistance requests shared in the interface. Like the participants, the autonomous robots also experienced the three information availability conditions. In Condition 1, the robots were only allowed to remember the location for two keys and two doors. This was to emulate how participants had to remember where they saw keys and doors because that information was not recorded on the minimap. In Condition 2, the

robots could remember the locations for all doors and keys they each saw, but they were not given access to any information other teammates’ gained, as this was the individual or personal mental model information condition. For the final condition, Condition 3, all teammates (robot and human) had access to the same information, as this shared set of information reflected the team’s common understanding or shared mental model.

4.6 Study Protocol

To begin participants reviewed an informed consent document, asked questions, and decided whether to participate, knowing they could stop at any time without penalty. After consenting, they watched an introductory video and completed a 15-minute tutorial on the search task, learning about the environment, human-agent interface, navigational controls, and key mechanisms. They also learned about the think-aloud method and real-time situation awareness questions. Before each task, participants watched a 30-second video about their assigned study condition. Each search task lasted up to eight minutes, followed by surveys assessing situation awareness and trust. At the end of the session, they answered demographic questions and provided feedback on their experience. The total session lasted around 90 minutes.

4.7 Study Measures and Metrics

How a participant relied on and trusted their autonomous teammates was of great interest in this study. During each simulation, participants were evaluated on their performance for the overall search task and reliance tasks. After completing the search task, the Human Robot Trust Model (HRTM) questionnaire^{1,2} measured the participant’s trust in each autonomous agent.

5. RESULTS

A mixed-methods analysis was performed to assess how the three different information availability conditions (C1 = No Information, C2 = Personal Information, C3 = Shared Information) influenced overall performance (duration, success rate), reliance task performance (teammate involvement, amount, success, trust outcome), and trust (evaluated by the HRTM). The Shapiro-Wilk test for normality concluded that duration, task success, and the HRTM data were not normally distributed. After performing multiple attempts to normalize the nonparametric data, the Friedman test, a nonparametric alternative to the repeated measures ANOVA, Kendall’s W coefficient, and pairwise testing using the Wilcoxon signed rank test with Bonferroni’s correction were used to analyze the nonparametric data. Kendall’s W coefficient, which follows Cohen’s interpretation of effect size,^{19,20} is considered a small effect from 0.1 to < 0.3, moderate effect from 0.3 to < 0.5, and a large effect if >= 0.5. Three participants were excluded from the data analysis for either failing to follow instructions or refusing to answer survey questions. Statistical testing did not show changes to the significance of results, having removed the data for these three participants. Table 1 displays the average values for the study measures across the three conditions (C1 = No Info, C2 = Personal Info, C3 = Shared Info).

Table 1. Mean values for team performance measures per information availability condition. Best scores shown in bold. Significant scores are enclosed in a rectangle.

Condition	Duration (min:sec)	Success Rate	HRTM for R1 (out of 5)	HRTM for R2 (out of 5)
No Info (C1)	8:00	3.51%	3.09	3.05
Personal Info (C2)	7:59	1.75%	3.01	3.05
Shared Info (C3)	7:11	38.60%	3.20	3.15

5.1 Search Task Performance

The time taken to complete the task was recorded in minutes, representing how long each participant took to reach the target within the maze environment, with a maximum limit of eight minutes. The Friedman rank sum test revealed a statistically significant difference in task duration across the three information availability conditions, $X^2(2) = 39.13$, $P < 0.0001$, with a moderate effect size ($W = 0.34$). Pairwise comparisons using the Wilcoxon signed-rank test showed significant differences in task duration between conditions C1 and C3 ($P < 0.0001$) and between C2 and C3 ($P = 0.0001$).

To complete the maze successfully, participants had to reach the hidden target within the 8-minute time limit. Even if an autonomous teammate located the target, participants were still required to reach it before the time limit. The Friedman rank sum test showed a statistically significant difference in task success across the three information availability conditions, $X^2(2) = 36.61$, $P < 0.0001$, with a moderate effect size ($W = 0.32$). Pairwise comparisons using the Wilcoxon signed-rank test revealed significant differences in participant success between Conditions 1 and 3 ($P < 0.0001$) and Conditions 2 and 3 ($P < 0.0001$).

5.2 Reliance Task Events and Resulting Grudge Levels

Because participants were not required to work with the autonomous teammates to complete the search task, relying or not relying on the agents for help finding keys was a choice made by the participants (and agents). This section analyzes the data regarding reliance and trust levels during the simulation and perceived trust in the autonomous agents post-task. Fig. 2 displays the average number of requests made as well as those accepted, ignored, and rejected. Highest mean values for each condition are shown in bold.

Table 2. Mean Number of Reliance Requests Among Teammates (Team Totals)

Condition	Statistic	Total Requests	Total Accepted Requests	Total Ignored Requests	Total Rejected Requests
No Info (C1)	avg (max/min) sd	14.65 (31/7) 5.03	7.04 (10/3) 1.60	12.44 (28/0) 18.81	0.19 (3/0) 0.55
Personal Info (C2)	avg (max/min) sd	17.28 (33/7) 5.78	6.33 (10/3) 1.52	14.12 (30/3) 8.98	0.33 (4/0) 0.87
Shared Info (C3)	avg (max/min) sd	7.74 (25/1) 4.50	3.63 (6/1) 1.65	5.25 (24/0) 4.24	0.09 (3/0) 0.43

5.2.1 Accepted Requests

Table 3 displays the average number of requests *accepted* by the team and each teammate per information availability condition. The **highest** values are shown in bold. In Condition 1 ($C1 = 7.04$) participants and autonomous teammates relied on each other for help finding keys most often, having the least amount of information. Condition 2 ($C2 = 6.33$) saw a decrease in the number of reliance tasks that took place, and Condition 3 ($C3 = 3.63$) showed the least number of reliance tasks between participants and autonomous teammates. The Friedman rank sum test indicated a statistically significant difference in the number of reliance tasks that occurred as information availability increased ($X^2(2) = 63.58$, $P < 0.0001$, $W = 0.56$ (large effect)). The pairwise Wilcoxon signed rank test found statistically significant differences between C1 and C2 ($P = 0.037$), between C2 and C3 ($P < 0.0001$), and between C1 and C3 ($P < 0.0001$).

5.2.2 Ignored Requests

Table 4 displays the average number of requests the team and each teammate *ignored*, including unanswered requests, per information availability condition. The **highest** values are shown in bold. In Condition 2 ($C2 = 14.12$) participants and autonomous teammates ignored each others' requests for help finding keys most often. In Condition 1 ($C1 = 12.44$) there was a decrease in the number of ignored reliance requests, and Condition 3 ($C3 = 5.25$) showed the least number of ignored reliance requests between participants and autonomous teammates. The Friedman rank sum test indicated a statistically significant difference in the number of ignored reliance

Table 3. Mean Number of Accepted Reliance Requests Between Participants (H) and Robots (R1, R2)

Condition	Statistic	Total Accepted Requests	# H's Requests Acc. by R1 R2	# R1's Requests Acc. by R2 H	# R2's Requests Acc. by R1 H
No Info (C1)	avg sd	7.04 1.60	0.44 0.42 0.80 0.78	2.58 0.75 1.41 1.06	2.39 0.46 1.79 0.83
Personal Info (C2)	avg sd	6.33 1.52	0.56 0.35 0.78 0.55	1.12 0.63 0.98 0.90	3.37 0.30 1.46 0.57
Shared Info (C3)	avg sd	3.63 1.65	0.86 0.19 1.03 0.44	1.09 0.16 1.09 0.53	1.19 0.14 1.11 0.35

requests as information availability increased ($X^2(2) = 43.88$, $P < 0.0001$, $W = 0.38$ (moderate effect)). The pairwise Wilcoxon signed rank test found statistically significant differences between C2 and C3 ($P < 0.0001$) and between C1 and C3 ($P < 0.0001$).

Table 4. Mean Number of Ignored Reliance Requests Between Participants (H) and Robots (R1, R2)

Condition	Statistic	Total Ignored Requests	# H's Requests Ig. by R1 R2	# R1's Requests Ig. by R2 H	# R2's Requests Ig. by R1 H
No Info (C1)	avg sd	12.44 18.81	2.37 2.28 2.49 2.35	2.70 0.04 2.67 0.19	4.89 0.16 18.72 0.62
Personal Info (C2)	avg sd	14.12 8.98	2.78 2.37 2.72 2.46	2.21 0.11 3.12 0.41	6.60 0.05 6.19 0.23
Shared Info (C3)	avg sd	5.25 4.24	1.56 1.37 1.61 1.58	1.35 0.02 1.68 0.13	0.93 0.02 1.51 0.13

5.2.3 Rejected Requests

Table 5 displays the average number of requests the team and each teammate *rejected* per information availability condition. The **highest** values are shown in bold. In Condition 2 ($C2 = 0.33$) participants and autonomous teammates rejected each others' requests for help finding keys most often. There were slightly less requests rejected in Condition 1 ($C1 = 0.19$), and Condition 3 ($C3 = 0.09$) resulted in the least number of rejected reliance requests between participants and autonomous teammates. The Friedman rank sum test did not indicate a statistically significant difference in the number of rejected reliance requests as information availability increased.

Table 5. Mean Number of Rejected Reliance Requests Between Participants (H) and Robots (R1, R2)

Condition	Statistic	Total Rejected Requests	# H's Requests Rej. by R1 R2	# R1's Requests Rej. by R2 H	# R2's Requests Rej. by R1 H
No Info (C1)	avg sd	0.19 0.55	0.00 0.00 0.00 0.00	0.04 0.02 0.26 0.13	0.05 0.09 0.29 0.29
Personal Info (C2)	avg sd	0.33 0.87	0.09 0.00 0.39 0.00	0.07 0.02 0.32 0.13	0.16 0.00 0.75 0.00
Shared Info (C3)	avg sd	0.09 0.43	0.02 0.00 0.13 0.00	0.00 0.00 0.00 0.00	0.05 0.02 0.40 0.13

5.2.4 Giving Up on Requests

Table 6 displays the average number of reliance tasks on which each robot *gave up*, per information availability condition. The **highest** values are shown in bold. In Condition 1 ($C1 = 2.16$) Robot 1 gave up on its teammates' requests for help finding keys most often. In Condition 2 ($C2 = 0.98$) there was a decrease in the number of reliance requests on which the Robot 1 gave up, and in Condition 3 ($C3 = 1.35$), Robot 1 gave up on the least

number of reliance tasks. The Friedman rank sum test indicated a statistically significant difference in the number of requests on which Robot 1 gave up as information availability increased ($X^2(2) = 21.69$, $P < 0.0001$, $W = 0.19$ (small effect)). The pairwise Wilcoxon signed rank test found statistically significant differences between C1 and C2 ($P < 0.0001$) and between C1 and C3 ($P = 0.007$).

Table 6. Mean Number of Reliance Tasks On Which Robots (R1, R2) Gave Up

Condition	Statistic	Robot 1	Robot 2
No Info (C1)	avg	2.16	1.93
	sd	1.24	1.44
Personal Info (C2)	avg	0.98	2.88
	sd	0.92	1.23
Shared Info (C3)	avg	1.35	0.95
	sd	1.19	1.16

In Condition 2 ($C2 = 2.88$) Robot 2 gave up on its teammates' requests for help finding keys most often. In Condition 1 ($C1 = 1.93$) there was a decrease in the number of reliance requests on which the Robot 2 gave up, and in Condition 3 ($C3 = 0.95$), Robot 2 gave up on the least number of reliance tasks. The Friedman rank sum test indicated a statistically significant difference in the number of requests on which Robot 2 gave up as information availability increased ($X^2(2) = 51.05$, $P < 0.0001$, $W = 0.45$ (moderate effect)). The pairwise Wilcoxon signed rank test found statistically significant differences between C1 and C2 ($P = 0.0001$), between C1 and C3 ($P = 0.0009$), and between C2 and C3 ($P < 0.0001$).

5.2.5 Completing Requests

Table 7 displays the average number of requests the autonomous teammates *completed* per information availability condition. The **highest** values are shown in bold. In Condition 2 ($C2 = 2.68$) Robot 1 successfully completed teammates' requests for help most often. In Condition 1 ($C1 = 1.09$) there was a decrease in the number of completed reliance tasks, and Condition 3 ($C3 = 0.86$) showed the least number of completed reliance tasks by Robot 1. The Friedman rank sum test indicated a statistically significant difference in the number of completed requests by Robot 1 as information availability increased ($X^2(2) = 52.51$, $P < 0.0001$, $W = 0.46$ (moderate effect)). The pairwise Wilcoxon signed rank test found statistically significant differences between C1 and C2 ($P < 0.0001$) and between C2 and C3 ($P < 0.0001$).

Table 7. Mean Number of Reliance Tasks Robots (R1, R2) Completed

Condition	Statistic	Robot 1	Robot 2
No Info (C1)	avg	1.09	0.88
	sd	1.06	1.02
Personal Info (C2)	avg	2.68	0.74
	sd	1.14	0.88
Shared Info (C3)	avg	0.86	0.81
	sd	1.04	0.85

In Condition 1 ($C1 = 0.88$) Robot 2 successfully completed teammates' requests for help most often. In Condition 2 ($C2 = 0.74$) there was a decrease in the number of completed reliance tasks, and Condition 3 ($C3 = 0.81$) showed the least number of completed reliance tasks by Robot 2. No statistical difference was found for the number of completed requests by Robot 2 as information availability increased.

5.2.6 Resulting Grudge Scores

Table 8 shows the average final grudge scores (and standard deviations) that resulted from the reciprocal cooperation, or lack thereof, between all teammates. The **lowest** values are shown in bold. The Friedman rank sum test indicated statistically significant differences in Robot 1's grudge against Robot 2 across the conditions

($X^2(2) = 15.27$, $P < 0.0005$, $W = 0.13$ (small effect)). The pairwise Wilcoxon signed rank test found statistically significant differences between C1 and C2 ($P = 0.019$) and between C2 and C3 ($P = 0.0002$). No statistically significant difference was seen in Robot 2’s grudge against Robot 1 across the conditions. Statistically significant differences were found for Robot 1’s grudge against the Human (participants) across the conditions ($X^2(2) = 40.03$, $P < 0.0001$, $W = 0.35$ (moderate effect)). The pairwise Wilcoxon signed rank test found statistically significant differences between C1 and C2 ($P < 0.0001$) and between C2 and C3 ($P < 0.0001$). No statistically significant difference was seen in Robot 2’s grudge against the Human (participants) across the conditions.

Table 8. Average Grudges Between Teammates (imp-un means grudge scores fall between implicit and untrust)

Condition	Measure	R1’s Grudge w/R2	R2’s Grudge w/R1	R1’s Grudge w/H	R2’s Grudge w/H
No Info (C1)	AVG SD	1.05 (imp-un) 1.17	1.21 (imp-un) 1.58	0.75 (implicit) 1.04	0.95 (implicit) 1.17
Personal Info (C2)	AVG SD	2.19 (untrust) 3.39	1.19 (imp-un) 1.25	2.68 (untrust) 3.62	0.89 (implicit) 1.30
Shared Info (C3)	AVG SD	0.67 (implicit) 0.69	0.91 (implicit) 1.06	0.47 (implicit) 0.97	0.63 (implicit) 0.96

5.3 Post-Task Perceived Trust

The HRTM was used to assess participants’ trust levels in each autonomous teammate, following the completion of a search task for each condition. Table 1 displays participants trust in Robot 1 and Robot 2 across the three information availability conditions. With an average trust value of 3.20 out of a possible 5 points, participants had the highest amount of trust for Robot 1 in Condition 3. They reported having the lowest trust in Robot 1 in Condition 2 with a slightly higher trust for Robot 1 in Condition 1. The Friedman rank sum test did not indicate a statistically significant difference for participants’ trust in Robot 1 or their trust in Robot 2 across the three information availability conditions. Participants’ trust in Robot 2 remained steady in Conditions 1 and 2 and improved in Condition 3, but participants did not seem to trust Robot 2 as much as they reported trusting Robot 1.

6. DISCUSSION

In an initial study, three levels of available information (no personal or shared information, personal but no shared information, and shared information) were provided to a non-hierarchical team completing a collaborative search task. It was expected to find that as more information became available, collaboration among teammates resulted in more efficient and improved performance as well as higher levels of trust. The results of the study do not directly reflect a gradual increase in trust and reliance as more information became available but, instead, show how partial information availability can significantly change trust and reliance among teammates. They also reflect that while collaboration, seen as the number of accepted reliance requests, decreased as more information became available, a team’s task duration and success rate improved.

6.1 Search Task Performance

For task duration, the results of the statistical tests performed during data analysis indicated statistically significant differences between Conditions 1 and 3 and between Conditions 2 and 3. No statistical difference was found between Conditions 1 and 2. For the successful completion of the maze, statistically significant differences were found between Conditions 1 and 3 and between Conditions 2 and 3. Again, no statistically significant difference was found between Conditions 1 and 2. Nevertheless, these results demonstrate that the overall performance of participants increases when the information available increases; however, the increase is not gradual from Condition 1 to Condition 2 to Condition 3. One possible explanation is that the eight-minute time constraint was too restricting, and with a small addition of time, more participants might have found the target in Condition 2. However, the teamwork reflected in Condition 2 may be interpreted as teammates using their personal

information of the environment for their own gain rather than to help their teammates, seeing as Condition 2 had the most ignored and rejected requests. Although there were only slightly less requests accepted by teammates in Condition 1 than in Condition 2, it appears that Robot 2 gave up on significantly more requests in Condition 2 than in Conditions 1 or 3. Unfortunately it is difficult to discern how often participants gave up on requests, as they did not always track (by hitting the cancel button) when they gave up on requests.

It was expected for Condition 3 to yield the best performance (duration and success) from participants; however, it was unexpected to see significantly fewer reliance requests in this condition. This could be that having all of the shared information available meant teammates did not feel that they needed to rely on each other as much and, instead, used the information provided to ‘clear the maze’ by opening as many doors as possible, thereby removing matching door/key pairs from the minimap. It could also be that participants viewed relying on each other to find and swap keys was no longer efficient because they could visually see all of the door and key locations that the teammates had found. They were no longer having to search for keys. It is likely that how teammates were relying on each other changed from searching for keys to bringing each other keys depending on the teammate closest to the key. This possibility was confirmed when reviewing comments made by participants during the search task.

6.2 Collaboration, Reliance, and Trust

Observations from data collection and analysis, combined with the results of the study, imply that the different information availability conditions influenced the way the teammates’ viewed teamwork and how they relied on another. For instance, the no information condition (C1) caused participants to rely on each other more heavily for any help they could get finding keys because they had very little information about the environment to guide them. The personal information (C2) caused the teammates to act more as individuals completing the same task while in a team but without relying on each other as much for help finding keys. Additionally, in Condition 2, the choice of which keys to ask for help finding changed and became more deliberately about the keys not reflected in each teammate’s personal information. In condition 3, however, having access to everyone’s shared information influenced teammates to ask for help retrieving keys rather than finding keys, making the team much more efficient and successful in completing the search task.

The results from the reliance task data analysis demonstrated statistically significant reductions in the number of reliance tasks that took place between teammates as the level of available information increased. Rather than this result being interpreted as teammates rejecting the optimal strategy of always asking for help when a key location was unknown, it is much more likely that as information availability increased, teammates did not need to ask for help as often because more information was being tracked (made available) on the minimap, and, therefore, more was known regarding key locations. Unfortunately, trust, both measured during the simulation and after, did not follow a similar or opposite trend, i.e. decreasing or increasing as information availability increased. Trust fluctuated from Condition 1 to Condition 2 and then Condition 2 to Condition 3 and was mostly the worst in Condition 2 for both measures. For the HRTM, participants’ trust was highest for both robots in Condition 3. With this condition showing and sharing the most information among all teammates, participants were able to visually see the robots ‘clear’ the maze of doors and keys as they moved through the environment, which might have led to increased trust. Additionally, the robots had more information available to guide them and keep them busy, potentially freeing participants from constantly having to help them and increasing participants’ trust in the robots. However, as shown in the data analysis performed on the HRTM results, no statistically significant difference was found for participants’ trust in either robot as the information availability increased.

Looking at the fluctuation in grudge scores among teammates (Table 8), it is interesting that Robot 1’s grudges against the Human (participants) and Robot 2 follow the same trend as the HRTM, but Robot 2’s grudges do not. In the simulations for the different conditions, there was always the possibility of one teammate being left out during reliance tasks due to the timing of the tasks and collaborative behavior of the robots. From observation during data collection, participants were usually the teammate excluded from reliance tasks because the robots were able to answer (and mostly accept) each other’s requests much faster than participants were able. This could also, in part, be why more reliance tasks occurred in Condition 1. With the number of reliance tasks that were ignored, rejected, or given up on, it is straightforward to see why Condition 2 produced worse grudge

scores than Conditions 1 or 3. Additionally, the grudge scores being similar (not significantly different) between Conditions 1 and 3 could be indicative of the teammates being more willing to work together, as they either had very minimal information to operate off of or almost too much information to guide them. Alternatively, Condition 3 could have the best grudge scores because it saw the fewest reliance tasks, and, therefore, there were less opportunities for teammates to ignore, reject, or give up on each other. However, it is much more likely that grudge scores are better in Condition 3 due to teammates’ access to the shared information.

One limitation that should be discussed is that, at the start in Condition 2, Robot 1 was blocked off from its other teammates. Participants and Robot 2 were able to get to each other to swap keys if performing a reliance task, but a door had to be opened before a key could be swapped with Robot 1. This was not a very limiting design constraint, as most of the participants unknowingly and very quickly unlocked the specific door blocking Robot 1. However, there were participants who noticed that Robot 1 was constrained to one side of the maze and felt that Robot 1 probably had access to the key to unlock the door. This was not the case, as the key for that particular door resided in the part of the maze where Robot 2 and the participants were traversing. Ultimately, this limitation was only limiting in performing reliance tasks with Robot 1 because the three keys required to reach the target were all accessible to the participants and Robot 2 from the start of the simulation.

7. CONCLUSION AND FUTURE WORK

In a study investigating human-agent teaming, information availability shaped teaming dynamics, trust, and reliance on teammates. In Condition 1 (no information), participants depended heavily on each other to locate keys due to limited guidance. In Condition 2 (personal information), they worked more independently, requesting help only for keys absent from their own recorded information. In Condition 3 (shared information), teammates focused on retrieving rather than finding keys, leading to greater efficiency and task success. In some cases, information availability had the unexpected effect of helping trust maintenance while reducing reliance on teammates. This was seen as participants mainly working by themselves, such that when information about key locations was available on the minimap, participants tended to not rely on the autonomous teammates to bring them keys but retrieved the keys themselves. In some instances, this had no impact on how quickly participants were able to retrieve a key and unlock a door, but there were cases where it would have been advantageous for participants to rely on their autonomous teammates, e.g., when the robots were closer to a key.

Future work will include further investigation of autonomous teammate behavior errors and if this trust-reliance system can identify and protect other autonomous teammates from falling prey to these errors. Additionally, further examination will be conducted to help define how information sharing and common understanding play a role in bridging the gap between trust and reliance, as these teaming dynamics play a vital role in integrating humans and autonomous agents into efficient teams.

ACKNOWLEDGMENTS

The authors wish to acknowledge the technical and financial support of the Automotive Research Center (ARC) in accordance with Cooperative Agreement W56HZV-24-2-0001 U.S. Army DEVCOM Ground Vehicle Systems Center (GVSC) Warren, MI. This material is based upon Cindy Bethel’s work supported while serving at the National Science Foundation. Any opinion, findings, and conclusions or recommendations expressed in this material are those of the author(s) and do not necessarily reflect the views of the National Science Foundation.

REFERENCES

- [1] Gulati, S., Sousa, S., and Lamas, D., “Design, development and evaluation of a human-computer trust scale,” *Behaviour and Information Technology* **38**(10), 1004–1015 (2019).
- [2] Simões, A. P. S. S. A. and Santos, J., “A trust scale for human-robot interaction: Translation, adaptation, and validation of a human computer trust scale,” *Human Behavior and Emerging Technologies* **2022**, 12 (2022).
- [3] Bohn, K. and Carter, G., “Reciprocal altruism and cooperation for mutual benefit,” in [*Encyclopedia of Evolutionary Psychological Science*], Shackelford, T. and Weekes-Shackelford, V., eds., Springer, Cham (2016).

- [4] Schweinfurth, M. K., “Cooperative intentions and their implications on reciprocal cooperation in norway rats,” *Ethology* **127**, 865–871 (09 2021).
- [5] Taborsky, M. and Riebli, T., “Coaction vs. reciprocal cooperation among unrelated individuals in social cichlids,” *Frontiers Ecology Evolution* **7** (01 2020).
- [6] Esterwood, C. and Robert, L. P., “The theory of mind and human–robot trust repair,” *Science Reports* **13**(9877) (2023).
- [7] Parasuraman, R. and Manzey, D. H., “Complacency and bias in human use of automation: An attentional integration,” *Human Factors* **52**(3), 381–410 (2010).
- [8] Lee, J. D. and See, K. A., “Trust in automation: Designing for appropriate reliance,” *Human Factors* **46**(1), 50–80 (2004).
- [9] Wang, N., Pynadath, D. V., and Hill, S. G., “Trust calibration within a human-robot team: Comparing automatically generated explanations,” in *[2016 11th ACM/IEEE International Conference on Human-Robot Interaction (HRI)]*, 109–116, Association for Computing Machinery (2016).
- [10] Lyons, J. B., Sadler, G. G., Koltai, K., and Battiste, H., “Shaping trust through transparent de-sign: Theoretical and experimental guidelines,” in *[Advances in Human Factors in Robots and Unmanned Systems. Advances in Intelligent Systems and Computing]*, Savage-Knepshield, P. and Chen, J., eds., **499**, 127–136, Springer, Cham (2017).
- [11] de Visser, E. J., Cohen, M., Freedy, A., and Parasuraman, R., “A design methodology for trust cue calibration in cognitive agents,” in *[Virtual, Augmented and Mixed Reality. Designing and Developing Virtual and Augmented Environments. VAMR 2014, Part I. Lecture Notes in Computer Science]*, Shumaker, R. and Lackey, S., eds., **8525**, 251–262, Springer, Cham (2014).
- [12] Helldin, T., Falkman, G., Riveiro, M., and Davidsson, S., “Presenting system uncertainty in automotive uis for supporting trust calibration in autonomous driving,” in *[AutomotiveUI '13: Proceedings of the 5th International Conference on Automotive User Interfaces and Interactive Vehicular Applications]*, 210–217 (2013).
- [13] Huang, S. H., Bhatia, K., Abbeel, P., and Dragan, A. D., “Establishing appropriate trust via critical states,” in *[2018 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)]*, 3929–3936, Association for Computing Machinery (2018).
- [14] Bindewald, J. M., Rusnock, C., and Miller, M. E., “Measuring human trust behavior in human-machine teams,” in *[Advances in Human Factors in Simulation and Modeling. AHFE 2017. Advances in Intelligent Systems and Computing]*, Cassenti, D., ed., **591**, ch. 12, Springer, Cham (2018).
- [15] Xu, A. and Dudek, G., “Optimo: Onling probabilistic trust inference model for asymmetric human-robot collaborations,” in *[10th ACM/IEEE International Conference on Human-Robot Interaction (HRI)]*, 221–228, Association for Computing Machinery (2015).
- [16] Okamura, K. and Yamada, S., “Calibrating trust in human-drone cooperative navigation,” in *[2020 29th IEEE International Conference on Robot and Human Interactive Communication (RO-MAN)]*, 1274–1279 (2020).
- [17] Ye, S., Neville, G., Schrum, M., Gombolay, M., Chernova, S., and Howard, A., “Human trust after robot mistakes: Study of the effects of different forms of robot communication,” in *[2019 28th IEEE International Conference on Robot and Human Interactive Communication (RO-MAN)]*, 1–7 (2019).
- [18] Mooers, B., Aldridge, A. L., Buck, A., Bethel, C. L., and Anderson, D. T., “Human-robot teaming for a cooperative game in a shared partially observable space,” in *[Proceedings of SPIE Defense + Commercial Sensing 12525, Geospatial Informatics XIII, 125250B]*, (2023).
- [19] Tomczak, M. and Tomczak-Lukaszewska, E., “The need to report effect size estimates revisited. an overview of some recommended measures of effect size,” *Trends in Sport Sciences* **01**(21), 19–25 (2014).
- [20] Cohen, J., *[Statistical power analysis for the behavioral sciences]*, Lawrence Erlbaum, Hillsdale, NJ, 2 ed. (1988).