

Design and Effects of Co-Learning in Human-AI Teams

van den Bosch, Karel; van Zoelen, Emma M.; Schoonderwoerd, Tjeerd A.J.; Solaki, Anthia; van der Stigchel, Birgit; Akrum, Ivana

DOI

[10.1613/jair.1.16846](https://doi.org/10.1613/jair.1.16846)

Publication date

2025

Document Version

Final published version

Published in

Journal of Artificial Intelligence Research

Citation (APA)

van den Bosch, K., van Zoelen, E. M., Schoonderwoerd, T. A. J., Solaki, A., van der Stigchel, B., & Akrum, I. (2025). Design and Effects of Co-Learning in Human-AI Teams. *Journal of Artificial Intelligence Research*, 82, 1445-1493. <https://doi.org/10.1613/jair.1.16846>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

Design and Effects of Co-Learning in Human-AI Teams

Karel van den Bosch

KAREL.VANDENBOSCH@TNO.NL

Human-Machine Teaming

Organization for Applied Scientific Research, TNO

Emma M. van Zoelen

E.M.VANZOELEN@TUDELFT.NL

Intelligent Systems Department

Delft University of Technology

Human-Machine Teaming

Organization for Applied Scientific Research, TNO

Tjeerd A.J. Schoonderwoerd

TJEERD.SCHOONDERWOERD@TNO.NL

Anthia Solaki

ANTHIA.SOLAKI@TNO.NL

Birgit van der Stigchel

BIRGIT.VANDERSTIGCHEL@TNO.NL

Ivana Akrum

IVANA.AKRUM@TNO.NL

Human-Machine Teaming

Organization for Applied Scientific Research, TNO

Abstract

The rapid progress of artificial intelligence (AI) will increase opportunities for humans and AI-driven technology to collaborate as teammates. This requires both partners to learn from interactions about the task, each other and the team (co-learning). Co-learning can be supported by enabling partners to share knowledge and experiences on the task and team level. This paper first analyzes the requirements regarding tasks and environments for co-learning. These requirements were subsequently implemented in a testbed: a human and intelligent robot jointly conducting an urban search and rescue task in a simplified task environment. We designed Learning Design Patterns (LDPs): interaction sequences intended to initiate and facilitate co-learning. Effects of LDPs on collaboration, knowledge and understanding, and team performance were experimentally evaluated using the testbed. In comparison to a previous study, participants appreciated the robot more, had more interaction and displayed more commitment. Results show evidence that the LDPs, in comparison with no interventions, initiated and improved learning of the human team member, in particular on knowledge development and understanding the partner. Better knowledge and understanding did, however, not also lead to better team performance. Implications for co-learning in human-AI teams and for learning-supporting interventions are discussed.

1. Introduction

Recent advances in Artificial Intelligence (AI) have been deployed in a variety of domains, such as healthcare (Sen, Kremer, & Buxbaum, 2019), finance (Gogas & Papadimitriou, 2021), transportation (Cunneen, Mullins, & Murphy, 2019) and urban search and rescue (Bartlett & Cooke, 2015; Demir, McNeese, & Cooke, 2018). AI may be used to control the behavior of devices that are by themselves not by definition intelligent, such as robots. By implementing algorithms that can tackle tasks requiring 'intelligence' (such as learning, perception, problem-solving, and logical reasoning), a robot with limited flexibility and adaptability can become an artificially intelligent robot, capable of (semi-)autonomously

performing complex tasks taking the context into account (Colagrossi, 2018). This paper uses the term AI to refer to AI-based intelligent systems having autonomy (O'Neill, McNeese, Barron, & Schelble, 2022), and a human-AI team as consisting of "one or more people and one or more AI systems requiring collaboration and coordination to achieve successful task completion" (National Academies of Sciences, 2021, p. 123) (Cuevas, Fiore, Caldwell, & Strater, 2007).

To make optimal use of the advances that AI has to offer, it is often necessary to go beyond the paradigm of human users operating systems as mere tools, to one whereby AI-based systems, such as intelligent agents and robots, interact with humans and work collaboratively towards a common goal (Bradshaw, Dignum, Jonker, & Sierhuis, 2012; Peeters, van Diggelen, van den Bosch, Bronkhorst, Neerinx, Schraagen, & Raaijmakers, 2021; Sycara & Sukthankar, 2006). There is an increasing interest in studying how humans and AI-based agents can effectively enhance their knowledge and capabilities both about their joint task and about each other (Ajoudani, Zanchettin, Ivaldi, Albu-Schäffer, Kosuge, & Khatib, 2018; Demir et al., 2018; O'Neill et al., 2022). This paper contributes to that field by discussing the conditions that bring about learning in human-AI teams, how these are implemented in a testbed for experimentation, and by presenting experimental research into the effects of learning interventions.

A prerequisite to any team undertaking coordinated action is establishing shared representations not only about the world but also about each other's representation of the world, a notion often referred to as *common ground* (Klein, Woods, Bradshaw, Hoffman, & Feltovich, 2004; Stalnaker, 2002), *common knowledge* (Fagin, Halpern, Moses, & Vardi, 1999), and *shared mental model* (Cannon-Bowers, Salas, & Converse, 1993; Mathieu, Heffner, Goodwin, Salas, & Cannon-Bowers, 2000; Jonker, Van Riemsdijk, Vermeulen, & Den Helder, 2010). This presupposes that team members keep internal representations (or "models") that contain knowledge about the ontic situation (e.g. of the joint task the team is working on), their own representations (e.g., own tasks, roles, capabilities) and of those of the other team members.

That is, it presupposes *Theory of Mind* (ToM), the cognitive capacity to attribute internal mental states like knowledge, beliefs, intentions, to oneself (self-awareness) but also to others (Premack & Woodruff, 1978). The term Theory of Mind is appropriate for humans but less so for AI-based agents, as it is hard to conceive them as having a 'mind'. However, to align with humans, AI-based agents too need a conceptualization that allows them to infer mental attitudes of others and to reason about these, because such agents are expected to operate in increasingly social and collaborative environments (da Silva, Rocha, Trajano, Morales, Sarkadi, & Panisson, 2024; Dunin-Keplicz & Verbrugge, 2011; Verbrugge, 2009). Implementations of ToM reasoning in AI-based agents and robots have thus attracted a lot of attention recently (Bolander, Dissing, & Herrmann, 2021; Devin & Alami, 2016; Dissing & Bolander, 2020; Li, Oguntola, Hughes, Lewis, & Sycara, 2022; Winfield, 2018). It is argued that AI-based agents equipped with a formal ToM of their human partner are more capable of inferring intentions and plans from interactions than agents that must induce such information from behavior observations alone. ToM simplifies the AI's learning about the human partner by making attributions based upon the model (Li et al., 2022). An implemented ToM enables the AI to behave and respond in alignment with the human, but it is probably different from how humans understand situations.

However, acquiring ToM (as humans) or implementing ToM (in AI-based systems) is only part of the story. Another challenge lies in *maintaining, refining, and using* shared representations across a dynamic context; this is far from trivial. Collaborative environments are often dynamic in nature: internal models need to be continuously adjusted in accordance with how events change the world (Van Den Bosch, Schoonderwoerd, Blankendaal, & Neerincx, 2019). Furthermore, applying a developed ToM to select behaviors that benefit the collaboration is not self-evident: it has been shown that it requires effort and experience to learn this in complex tasks (Apperly, Riggs, Simpson, Chiavarino, & Samson, 2006; Keysar, Lin, & Barr, 2003; Verbrugge & Mol, 2008). This paper discusses and investigates how learning in human-AI teams can be facilitated.

We use the term Human-AI *co-learning* (in this paper also abbreviated as co-learning) to define the process in which collaborating humans and AI adapt to each other and learn together over time (Holstein, Aleven, & Rummel, 2020; Van Zoelen, Van Den Bosch, & Neerincx, 2021a). Co-learning results from interactions between collaborating team members. Team members may recognize behavioral adaptations as successful, and subsequently apply these patterns of successful collaboration in new but similar situations. Co-learning can improve the performance of the human-AI team (similar to *mutual adaptation* in human-robot teaming literature (e.g., Nikolaidis, Hsu, & Srinivasa, 2017; Van Zoelen, Van Den Bosch, & Neerincx, 2021b)). In addition, co-learning may also be instigated by situations deliberately designed to elicit the interactions that are expected to bring about the learning that generalizes to new situations. Our notion of co-learning, as illustrated in Figure 1, emphasizes that both partners learn over time from their interactions: they learn about the task, about themselves, about their partner(s), and about how to distribute the tasks (Van Den Bosch et al., 2019).

In this paper we address the question how to instigate co-learning in a human-AI team by building on the notion of Design Patterns: generic, reusable, and proven solutions to a commonly occurring problem within a given context (Van Diggelen, Neerincx, Peeters, & Schraagen, 2018). *Learning* Design Patterns (LDPs) refer to activities and interactions intended to support members of a human-AI team to improve their mental models, which is envisioned to benefit team performance on future task situations. Schoonderwoerd and colleagues (Schoonderwoerd, Van Zoelen, Van Den Bosch, & Neerincx, 2022) designed and tested the effects of LDPs on learning and performance of a team consisting of a human and an intelligent robot when jointly conducting urban search and rescue (USAR) in a simulated task environment. They found LDPs to support humans’ knowledge, understanding and awareness of their robot partner. However, the improved understanding of the partner did not result in a more fluent collaboration, nor in improved performance of the team. It was concluded that humans are not naturally inclined to collaborate with intelligent agents as a team. Instead, humans tended to treat the AI-based robot as a tool rather than as partner. The authors observed that participants did not realize that engaging in learning activities would benefit the functioning of the team as whole. Suggestions are to make humans more aware that collaboration, alignment, and learning with an AI-based partner is beneficial to the team, and to expressively guide both partners in the activities that will bring this about (Schoonderwoerd et al., 2022). In section 2 we discuss the objective of supporting human team members to consider and treat AI-based agents as real team partners; in section 3 we present how the principles were implemented in a testbed for co-learning. Section 4 presents an experimental study into the effects of learning interventions. Sections 5 and 6 present discussion and conclusions.

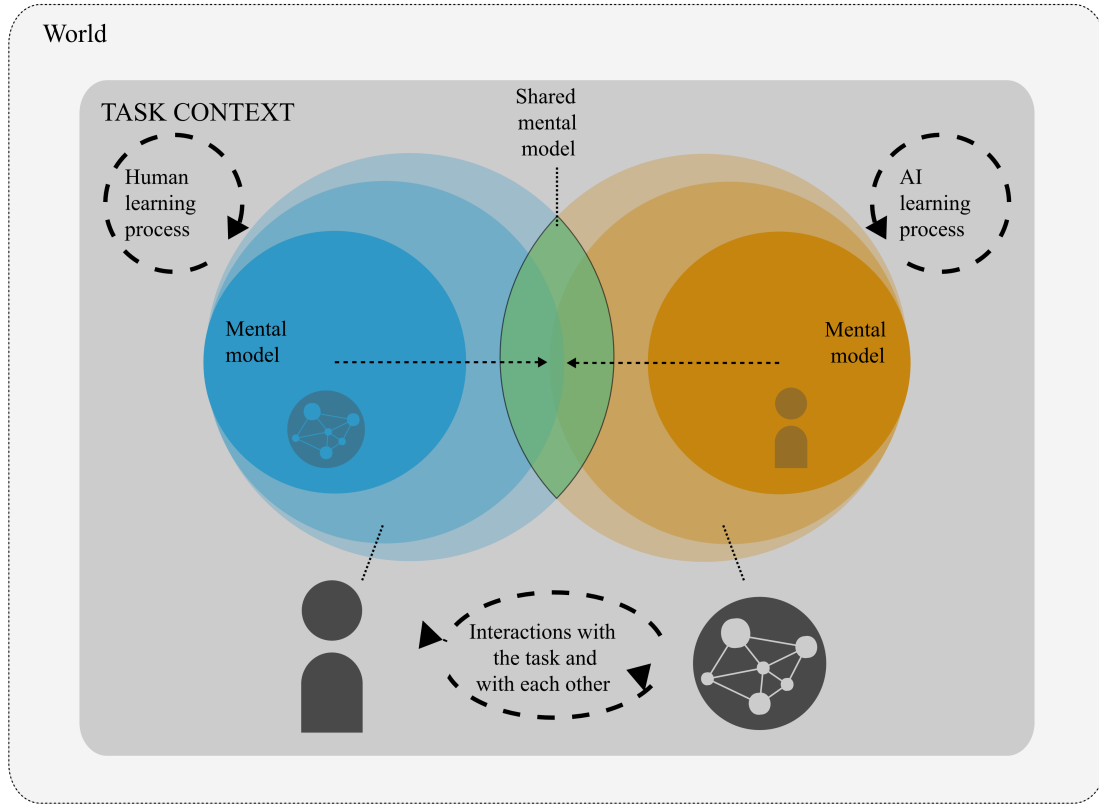


Figure 1: Human-AI co-learning: both partners learn over time from their interactions: they learn about the task, about themselves, about their partner(s), and about how to distribute the tasks. This learning progressively accumulates in a shared mental model.

1.1 The Present Study

The present study involves a theoretical analysis on the study of co-learning in human-AI teams, and an experiment with human participants to empirically investigate the effects of interventions (LDPs) on learning. Although co-learning evidently involves both partners simultaneously, in this study we designed learning interventions for studying how *humans* learn from co-learning interactions with an intelligent AI-based robot ¹. The theoretical analysis is on determining the requirements for enabling co-learning. These requirements pertain to the design of a task, of the task environment, and on the functionalities and behavior of the AI-driven agents. It also pertains to qualities of the human team member (e.g., participant). Building on work of others on this topic (e.g., Johnson, Bradshaw, Feltovich, Jonker, Van Riemsdijk, & Sierhuis, 2014; Johnson & Bradshaw, 2021; Schoonderwoerd et al., 2022; Van Zoelen et al., 2021a) we address the research questions:

RQ1: What are the requirements for enabling the study of co-learning in human-AI teams?

1. This is the first step towards the study of bi-directional learning. In the future we intend to encompass the AI's learning as well.

The requirements obtained by means of a theoretical analysis (see section 2) are subsequently used to develop a testbed (see section 3) suitable for empirical research. The testbed is a simulation of an *urban search and rescue* (USAR) task, in which a human participant and an autonomous intelligent robot jointly search for victims in an earthquake-area, and evacuate them to a command post where they are safe. Learning challenges for both the human and the robot are deliberately designed into the testbed.

Learning interventions, in the form of Learning Design Patterns have been developed, intended to support the team with achieving the learning challenges. We address the research question:

RQ2: What are the effects of Learning Design Patterns on co-learning within a human-AI team? In particular:

- What are the effects of LDPs on the participants' collaboration with the robot?
- What are the effects of LDPs on the participants' knowledge and understanding?
- What are the effects of LDPs on team performance?

2. Requirements for Studying Co-learning in Human-AI Teams

Our research concerns how humans and AI-based agents learn together, and how this learning may be intentionally supported by designed interventions. To study how processes of co-learning develop, and how they might be facilitated, requires a suitable testbed. The research field Human-Machine Teaming has produced various testbeds (see for example: (Barber, Davis, Nicholson, Finkelstein, & Chen, 2008; Demir, McNeese, Johnson, Gorman, Grimm, & Cooke, 2019; Harbers, Bradshaw, Johnson, Feltovich, Van Den Bosch, & Meyer, 2011; Holec, Hockensmith, Broll, Wittich, Donadio, de Visser, & Tossell, 2020; Johnson, Bradshaw, Duran, Vignati, Feltovich, Jonker, & van Riemsdijk, 2015; Lematta, Coleman, Bhatti, Chiou, McNeese, Demir, & Cooke, 2019; McNeese, Demir, Cooke, & She, 2021). See (Lematta et al., 2019) for a clear overview and discussion.

Based upon the literature, we determined the following requirements for a testbed to study co-learning:

1. It should support the design of dependency between tasks, and interdependency between humans and AI-based agents, which bring about the need for partners to collaborate;
2. It should provide ample opportunities for partners to adapt their behavior to actions of the partner, so that they can learn from the experiences and feedback;
3. It should provide opportunity for partners to reflect on experiences so that they can develop and refine their understanding of their own role in the team, of that of their team partner(s), and of how the team should organize the work to achieve the team's mission effectively and efficiently.

In the following sections, we discuss the theory and rationale underlying the requirements. We conclude with practical implications for research.

2.1 Dependency, Interdependency, and the Need to Collaborate

Co-learning in human-AI teams emanates from interactions during collaboration. A task suitable for studying co-learning should therefore support interaction. When partners'

capabilities are, at least to some extent, complementary with respect to these tasks, then interdependency between team members arises. For example, there is interdependency when a member lacks a required capability to competently perform a task in a given context and, as a result, has to rely on a partner to do the task. Interdependency may also arise when a partner has to regulate its behavior in response to the behavior of a team member that works on a related task, like when team success requires that partners synchronize their actions, or when a team member first needs to obtain permission from a partner before commencing its own task.

Interdependency implies that partners have to adapt their behavior to other team-members, for example by synchronizing, delegating, taking over, permitting or prohibiting actions, and so on (Johnson et al., 2014; Johnson & Vera, 2019; Johnson & Bradshaw, 2021). Interdependency can demand collaboration, for example when one team member cannot perform a particular task without the help of another team member. This is called a *hard* interdependency. Interdependency may also encourage collaboration. This may occur for example when a partner, although capable of accomplishing the task alone, can complete a task much better and quicker when assisted by a team mate. The latter type is called *soft* interdependency. Previous studies have shown that hard dependency is in itself not sufficient for the human to regard the agent as a team member (Schoonderwoerd et al., 2022).

2.2 Opportunities for Adaptation and Learning

To enable investigation of co-learning in a human-AI team, it is evidently essential that the study environment provides opportunities for adaptation and learning by both partners. This may be achieved by deliberately introducing knowledge gaps in either or both partners, which will bring about situations in which the team has to learn in order to solve the problem at hand more efficiently (soft dependency), or even to solve it at all (hard dependency). The challenge for the team is to first recognize that a problem or obstacle exists, and secondly, to determine and initiate activities producing the adaptation and learning that solves or alleviates the problem. Incorporating soft task interdependencies in the study environment ensures that team members' tasks are intertwined. As a result, the strategy of executing tasks individually is less efficient than when doing it in collaboration. Previous work has shown that partners generally do not learn this at the first opportunity, but that they require multiple interactions to commence adaptation and learning (Van Zoelen et al., 2021a, 2021b). A study environment therefore needs to support learning as a continuous and bidirectional process.

Two types of learning may be distinguished in a human-AI collaboration context: *emergent learning* and *intentional learning*. Emergent learning occurs while acting or interacting and in response spontaneously develop an understanding of the task and team (Critchfield & Twyman, 2014), offering partners a basis to adapt their behavior to the other (Van Zoelen et al., 2021a), which may either lead to successful (or less successful) team performance. For example, a team member observes that its partner struggles to solve a task on its own, and therefore decides to interrupt its own work to help the teammate. The partner receiving the assistance might learn that it can count on help from this teammate when it comes to solving this particular task. Successful interactions can be selected and sustained by partners and consolidated into a first version of a collaboration pattern (Van Zoelen et al., 2021a). When the team applies a collaboration pattern on subsequent occasions, the team uses the

feedback from the task environment to correct and refine the collaboration pattern. Partners may be aware of when they engage in emergent adaptation and learning, but these processes may also take place implicitly and unconsciously (Patterson, Pierce, Bell, & Klein, 2010). Experiences of performing actions and interactions with team mates in the task context are necessary to feed emergent learning (Critchfield & Twyman, 2014). The own and partner's behavior, along with their consequences in the task environment, produce the behavioral cues that form the input for emergent learning.

Another approach to learning in a human-AI team is intentional learning. In contrast to emergent learning, where the opportunities for the team to learn arise from their fortuitous interactions, intentional learning occurs in situations specifically designed to elicit the activities of the team that should bring about effective and efficient learning. For example, in a situation where a team mate cannot perform its task due to lacking knowledge, the team mate is requested to explain its performance to his team members, who may thus become aware of the knowledge deficiency. Engaging in predetermined learning activities enable team members to improve, correct and extend their mental models. It is believed that making team members explicitly aware of what has been learned supports them to sustain successful interactions beyond the training context. One way to design intentional learning is by means of Learning Design Patterns (LDPs) (Schoonderwoerd et al., 2022), that specify what activities need to be initiated by which partner in what situational context.

2.3 Reflection to Improve Mental Models

It is argued that when a team faces an obstacle or problem, partners seek to adapt their behavior in order to improve the team's performance. This adaptation may (temporarily) help, but does not directly lead to learning as not all relevant information is available at that moment. Learning requires partners to reflect upon how situations, tasks, capabilities, and collaboration relate to each other. Once understanding through reflection develops, and team members refine their mental models, they need to be able to share with their partners what they have learned (e.g., by explaining, by showing, by teaching). By sharing, communicating and reflection, partners can extend their model with the perspective of the partner in terms of its goals, values, needs, capabilities, resources, plans, emotions, and possibly other characteristics. In other words: the forming of a *Theory of Mind* (Van Den Bosch et al., 2019). Although humans have different learning processes and representations, their reflecting upon experiences is needed to attain an aligned understanding of their shared world. Such learning supports the team to perform tasks in an increasingly adaptive and effective manner.

In summary, a study environment for investigating human-AI learning should provide partners with ample opportunities to develop and tune their models. The task should induce team members to develop an accurate understanding of the world, of themselves, and of their teammates. Furthermore, the study environment should facilitate the processes that are needed for team partners to elaborate, correct and refine their models.

2.4 Practical Considerations for a Research Environment

Finally, aside from the requirements in the sections 2.1-2.3, experience from previous research has shown that the task- and workload of the human should be maintained at a moderate level. Because when loads are high, the tendency of humans is to 'lock-in' on their task, and to not pay sufficient attention to the events in the environment, and to the behavior of

the partner (Neerincx, 2007). High task- and work loads therefore create conditions that interfere with the opportunity to learn. The task’s complexity should be controlled at a level that allow humans with sufficient time and resources to learn.

3. Design and Implementation of Testbed

The USAR domain has been proposed as suitable for the design of testbeds that adhere to the principles mentioned in section 2 (Lematta et al., 2019; Schoonderwoerd et al., 2022). This is especially the case since AI-based robots are expected to increasingly take on cognitive and physical functions in response to disaster events (Liu & Nejat, 2013). Following this, we designed a virtual USAR task for a human-robot team (see section 3.1). In the task, several learning situations are introduced (see section 3.2). The behavior of the robot, and the conditions that must be met in order for the robot to learn during the task, were pre-defined and programmed into the robot (see section 3.3). Learning design patterns (LDPs) are provided to support the co-learning that the team needs to achieve the learning opportunities that are presented in the task (see section 3.4).

3.1 Description of the Task

Figure 1 shows a screenshot of the implemented USAR task. The narrative is: an earthquake has struck a village and it is likely that there are multiple victims inside the buildings. A human-robot rescue team has to explore all buildings and to rescue victims by extracting them from the disaster area as quickly as possible. The avatar of the human rescue worker is controlled by the participant, while the avatar of the robot is controlled by an AI-driven agent. Henceforth we use the term ‘human’ to refer to the participant and its avatar, and ‘robot’ to refer to the robot-avatar driven by the AI-agent. The mission of the rescue team is to find and rescue all victims by entering the red-bricked buildings and carrying the victims to the green-bricked command post. The human navigates through the environment by using the arrow keys and can select actions from a list (e.g., carry a victim) by right clicking the object of interest. The human can also send pre-defined messages to the robot in this way, for example a request to the robot to carry a particular victim. Communication is bidirectional, as the robot will respond to messages from the human, and also sends pre-defined requests and provides information by sending messages to the human. All communications are listed in Appendix A, and they appear in a chat box (see right panel of Figure 2).

There are five subtasks that, if performed correctly, contribute to the success of the team: *navigating*, *clearing building entrances*, *entering buildings*, *carrying victims*, and *taking shelter*.

Navigating

Team members can move up, down, left, and right within the tile grid representing the disaster area. They cannot move through walls, so a building can only be entered through its door. All ground tiles can be stepped on, but some tiles have become covered by mud as a result of the earthquake (mud is visualized as brown stains, see Figure 2). Mud significantly slows down the movement speed of a team member that traverses it. Delay in movement can be avoided by choosing routes that avoid mud-covered tiles.



Figure 2: Screenshot of the USAR task in which a participant-controlled human (avatar with yellow helmet) and AI-controlled robot (avatar with blue eye) are missioned to rescue victims from intact and collapsed buildings (red-bricked buildings) and bring them to the command post (green-bricked building). Victims can either be severely wounded (red) or slightly wounded (yellow). The area also contains mud pools (brown stains) that, when traversed, slows down movement. Communication between team members is shown in a chat box (shown on the right).

Clearing building entrances

Buildings are either intact (undamaged walls) or collapsed (damaged walls). The door of a collapsed building is inoperable and blocks entry to the building. The door therefore needs to be removed. Removing the door from the entrance can only be performed by joint action of the human and robot.

Carrying victims

Victims are slightly wounded (yellow icon) or severely wounded (red icon) (see Figure 2). Prior to the task, it has been pointed out to participants that the health condition of victims deteriorate over time while not yet in the command post. The purpose of this instruction was to urge the rescue team to find and rescue all victims as quickly as possible (although health deterioration was not actually implemented in the experiment). Slightly wounded victims can be carried by a single rescue team member, while severely wounded victims require joint action, as they have to be carried on a stretcher. Saving severely wounded victims from a collapsed building forms a challenge, as only the robot can enter a collapsed building. In these circumstances the robot can drag (not carry) the victim outside of the collapsed building. Subsequently the victim can be placed on a stretcher and carried by human and robot together.

Taking shelter

At certain predetermined moments during the mission, the area is hit by another earthquake, which imposes a threat to the rescue worker and the robot. The impact of a new earthquake on the rescue mission is implemented by slowing the movement speed of both robot and human, that persists for a certain amount of time. However, partners can prevent this negative effect by timely taking shelter in the doorway of a building. If a team member does not timely take shelter, their movement speed is temporarily delayed. This temporary effect is indicated with a red cross that shows on the avatar while the effect applies.

The task was designed to be simple to understand and to perform, leaving sufficient opportunities for the participant to build a mental model about the robot and to identify and solve knowledge gaps through communication with the robot. These knowledge gaps, learning situations, and potential interventions that occur during the task are further explained in section 3.3.

3.2 Description of Learning Opportunities

The learning opportunities directly follow from the interdependencies between subtasks. These were intentionally created by different informational deficits to the robot and the human. For example, only the robot knows that both actors are required to open broken doors, while only the human knows that they are both required to carry severely wounded victims. At the start of the mission, actors are unaware about their own knowledge gaps and those of the other team member. We thus created a situation in which team members had to learn about the abilities and shortfalls of the other team member.

Development of an AI-based system that learns through interaction with its environment and team members is very complex and requires prolonged exposure to learning situations over time. For this experiment we therefore chose to design specific learning situations for both team members and to predefine the conditions for learning to occur in the robot. The responses of the robot to interactions that were initiated by the human were also pre-specified. That is, when the robot encountered a learning opportunity during the task, it only successfully learned if the human showed teaching behavior such as cueing or instructing the robot to do something (for details see section 3.4). Evidently, whether or not the human or robot learned during the experiment depended upon their response to the encountered learning opportunities.

We developed two types of learning opportunities in the search-and-rescue task. In the first type, which we call *knowledge gap*, one team member lacks information that is required to perform a particular task and has to learn this from its partner, while the partner needs to learn that this team member lacks the task-critical knowledge. In the second type of learning opportunity, *joint task*, one team member has to learn that a task needs collaborative action, while its partner has to learn that this information is not known by this team member. Both types of learning opportunities were applied symmetrically in the testbed, resulting in four learning situations: one knowledge gap and one joint task opportunity in which the human has to learn about the task from the robot, and one knowledge gap and one joint task opportunity in which the robot has to learn about the task from the human.

These learning opportunities were implemented by introducing the following events in the scenario: (a) a second earthquake hitting the disaster area (*knowledge gap*); (b) mud covered paths (*knowledge gap*); (c) jammed doors of damaged buildings (*joint task*); and (d)

severely wounded victims requiring collaborative (rather than individual) carrying (*joint task*). These events created learning opportunities because crucial information was deliberately *not* provided to either the human or the robot. For the latter, 'not providing crucial information' implied that the robot was programmed to act without taking this knowledge into account. The learning challenges and potential learning outcomes are shown in Figure 3.

3.3 Description of Robot's Behavior

Figure 4 illustrates how the robot was programmed to choose its behavior during the task. The default behavior of the robot is to plan a path towards buildings that have not yet been inspected, to search for victims in these buildings, and to bring them to the command post. This continues until all victims are saved. However, the robot will come across four types of learning situations that complicate task execution: (1) mud paths, which slow down movement when traversed; (2) earthquakes, which slow down movement when not taking shelter; (3) severely wounded victims, who can only be carried by joint action; and (4) jammed doors, which can only be opened by joint action. Initially, the robot knows about (2) and (4), but does not know about (1) and (3). The latter two are however known by the participant and can also be learned by the robot through learning activities if these are initiated by the participant or the headquarters. From the moment that this information is learned by the robot, it will change its behavior by planning paths around mud and by helping the participant to carry severely wounded victims.

As can be seen in Figure 4, the decision flow of the robot starts with planning a path to any uninspected building that is left. This path is planned either through or around mud, depending on the current knowledge of the robot. Then, the robot checks whether any of the learning situations (2), (3) or (4) apply. The robot behavior in these events is explained below.

Taking shelter for an earthquake

Earthquakes occurring during the rescue operation are triggered after delivering the second-saved victim in the command post, and after delivering the fourth-saved victim. Shortly prior to an earthquake the robot will seek shelter in the nearest building entrance. This shelter-seeking acts as behavioral cue for the participant, who initially does not know that this behavior of the robot forecasts an earthquake.

If despite the robot's behavioral cueing the human does not seek shelter itself, then participants of the intentional learning group receive additional cues. These can have the form of: an explanation, an assignment, or an appeal to reflect. First, the robot will send a chat message in which it explains its behavior. If the participant is still not moving towards an entrance to seek shelter, the robot sends a chat message assigning the participant to take shelter. When the participant still does not respond adequately, the mission headquarters sends a message intended to trigger the participant to reflect on the behavior of the robot ('*We see that the robot is standing still in the doorway. Why is it doing that?*').

Carrying a severely wounded victim

If the earthquake learning situation does not apply, the robot checks whether the participant has requested the robot to assist with carrying a severely wounded victim. Participants in the *emergent learning condition* can only implicitly request help from the robot through a behavioral cue (i.e., by standing still for some time next to the victim). Participants in

Situation	Learning opportunities	Type of opportunity	Prior knowledge and capabilities	Potential learning outcomes **
Earthquake	(1) knowing when an earthquake is imminent	Knowledge gap	Robot has seismograph predicting imminent earthquakes;	Robot learns that human explorer does not know when earthquake is imminent Robot learns that human explorer does not know how to shelter
	(2) knowing when and where to take shelter		Participant is not instructed about possibility of new earthquakes, nor that robot has seismograph.	Participant learns that robot has tool to predict earth quake (r, <u>hq</u>)*** Participant learns to copy shelter behavior from robot (r, <u>hq</u>)
Mud on path	(1) being aware that mud on path slows movement	Knowledge gap	Participant has the opportunity and capability to observe that mud slows movement;	Participant learns that mud on paths slows down movement Participant learns that robot is unaware of slowing effect of mud (h, <u>hq</u>) Participant teaches robot to avoid mud (h)
	(2) navigating with minimal delay		Robot is pre-programmed not to know this.	Robot learns how to prevent delays when navigating
Jammed doors	(1) knowing that opening jammed doors need collaborative action	Joint action	Robot knows that opening jammed doors need joint action; Participant is instructed that robot opens doors, but is not informed that opening jammed doors need joint action.	Robot learns that human explorer is unaware that jammed doors need to be opened collaboratively Participant learns to jointly open jammed doors with robot (r, <u>hq</u>)
Severely wounded victims	(1) knowing that severely wounded victims need collaborative carrying	Joint action	Participant is instructed that light-wounded can be carried individually, but that severely wounded victims need to be carried together; Robot does not know this; Participant is not informed that robot does not know this.	Participant learns that robot is unaware that severely wounded victims need to be carried collaboratively (h, <u>hq</u>) Robot learns to jointly carry severely wounded victims with human explorer

Figure 3: Learning opportunities, created by events in the scenario. Team members differ in their prior knowledge and capabilities for each situation, but have the potential to learn from each other to successfully overcome the learning challenges.

* Co-learning involves learning by both partners. To emphasize the mutuality and complementarity of co-learning, we included learning outcomes for both the human explorer and the robot in the final column. However, as explained earlier, in this study we focus on learning by the human partner only. The behavior and the conditions for the learning by the robot are pre-programmed.

** the actor that performs the activities defined in the LDP can either be the human (h), the robot (r), and/or head quarters (hq). In brackets are shown which actors(s) perform the LDP activities for that particular learning outcome.

the *intentional learning condition* can also explicitly request help by explaining to the robot that severely wounded victims cannot be carried alone, or by assigning the robot to assist with carrying. If the participant requested help, the robot suspends its current action and

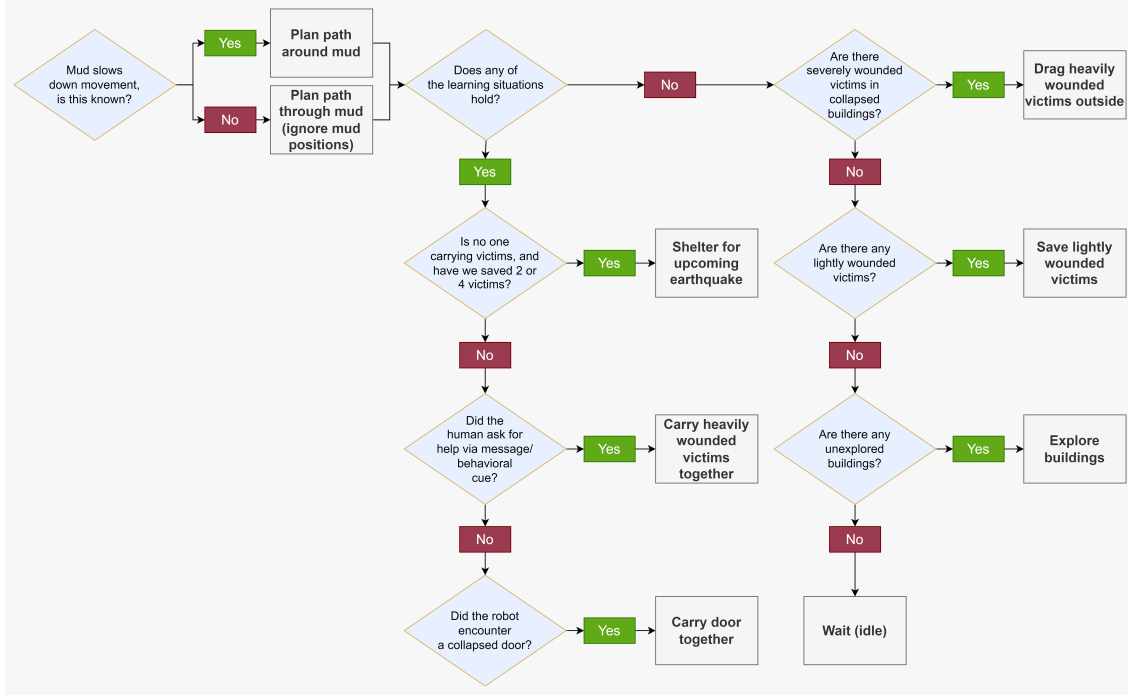


Figure 4: Decision diagram of the robot’s behavior in the search-and-rescue task.

immediately plans a path towards the victim. If the participant is still near the victim when the robot arrives, they will automatically pick up the victim together and their avatars will merge into one joint avatar, indicating that now both actors are carrying the victim. At that moment, the participant takes full control over the movements of the team. Once the victim is safely in the command post, the human and robot regain their own avatar again. After collaboratively saving one victim, the robot has automatically learned that joint action is required to carry severely wounded victims. Having learned this, the robot will autonomously move towards the human when (s)he is standing next to a severely wounded victim for a short amount of time.

Dislodging a collapsed door

If both the *earthquake* and *carry victim* learning scenarios do not apply, the robot checks whether it is standing next to a collapsed door. If so, it starts the *carry door together* sequence. This sequence starts with the robot sending a behavior cue to the participant by moving back and forth in front of the door. If necessary, participants of the *intentional learning group* receive additional cues, again in the form of an explanation, an assignment, or an appeal to reflect. If despite the robots’ cues, the participant is not moving towards the robot for some amount of time, the robot will first send a chat message with an explanation of its behavior. If the participant is still not responding, the robot sends an assignment via the chat to request assistance from the participant. If this also proves to be ineffective, headquarters will send a final message aimed to trigger reflection on the part of the participant (‘*We see that the robot is not coming to help you. Why would that be?*’). If the participant does not respond, the robot continues its behavior flow by planning a path to another building. If the participant does respond by moving towards the collapsed door, the robot will wait

until the participant arrives. When arriving, an animation starts showing the robot and the participant lifting the collapsed door and putting it aside. When this is finished, the robot enters the collapsed building to search for victims.

If none of the above learning situations apply, the robot first checks whether there are any severely wounded victims left in collapsed buildings. If so, it will prioritize dragging these victims outside the building so that they can be carried jointly with the participant to the command post. If there are no known severely wounded victims left, the robot will carry any known slightly wounded victims. If all known slightly wounded victims are in the command post, the robot will search for any unexplored buildings.

When a human-robot team is engaged in tasks that require interdependent actions, communication is very important (Lematta et al., 2019). In our task we enable bidirectional communication between participant and robot using a chat window. At the start of each scenario, the robot sends a message about the first two buildings to search (*'I will start at the Post Office, can you start at Dorpsstraat 1?'*). Additionally, each learning intervention involves sending and receiving messages between participant and robot in order to share knowledge and expectations that can help to successfully overcome the learning challenges.

All learning situations include timers and timing thresholds for the robot and the participant for behavioral cueing, for acquiring new knowledge, and for responding to learning interventions. For example, after the robot has sent an explanation why it stalls at an entrance, the robot will wait for a while to allow the participant to respond by also taking shelter. If the participant does not respond and the set time threshold is reached, the robot sends the next message (assignment). All timing thresholds were tuned on the basis of a pilot experiment that was conducted among a small group of participants (see Appendix E).

3.4 Description of Learning Activities

There are four situations presented in the task that present opportunities for team members to improve their performance by engaging in learning activities (see Figure 4). For each situation, one of the two team members lacks the required knowledge to effectively resolve the situation. Therefore, the other team member can initiate a learning activity to teach its partner the required knowledge. We implemented four of such learning activities in the task: a *behavioral cue*, allowing emergent learning to occur in the task, and an *explanation*, *assignment*, and *reflection*, intended to bring about intentional learning in the task (see section 3.2, for a description of emergent and intentional learning). Table 1 summarizes the developed LDPs and their use in the present study. The three intentional learning activities were implemented as Learning Design Patterns: activities that may be initiated in each learning situation by the human, robot, or headquarter agent, with the aim of teaching a team member about the task. More specifically, the explanation and assignment LDPs could be initiated by the human or the robot, while the reflection LDP could only be initiated by the headquarters. In this study we are interested in the effects of intentional learning in human-AI teams on performance, learning and collaboration. Therefore, we regard the behavioral cue as the reference standard, and investigate whether additional learning improvements can be achieved by implementing LDPs in the task.

3.4.1 BEHAVIORAL CUEING

When in the scenario a situation arises for the participant to learn (see Figure 4) then the robot first gives a behavioral cue, intended as a signal for the human. For example, the robot starts moving back and forth in front of a jammed door, thus giving the human the opportunity to observe that the robot does not open the door (as it does with non-jammed doors), and the opportunity to infer that something must be figured out to open the jammed door. It requires the human to deduce from observation, and to infer what solution might work. When a situation arises for the robot to learn the human may provide a behavioral cue to the robot. For example, the human might stand still next to a severely wounded victim for a while, thereby cueing the robot to come and help carry this victim. If this occurred (as detected by location tracking and timers), the robot was programmed to automatically notice the cue, and to respond by navigating towards the severely wounded victim.

3.4.2 INTENTIONAL LEARNING

We designed three types of intentional learning activities (LDPs) that aim to facilitate co-learning in a team: *Explain*, *Assign*, and *Reflect*. These activities are integrated into the task context and can be initiated in all four learning situations as described in Table 1. The LDPs contain activities intended to bring about learning by the human and by the robot. When a situation arises for the robot to learn, the human may issue an explanation or assignment to the robot using a context-menu of the interface. An example of an explanation from the human to the robot is: '*Heavily wounded victims need to be carried together!*'. When a situation arises for the human to learn, and the human does not respond to the robot's behavioral cues, then the robot issues an explanation or assignment (e.g., the assignment '*Can you walk to the nearest doorway*'). An explanation aims to make the partner aware *why* particular behavior is needed, while an assignment provides a behavior instruction to the recipient. The third LDP (*Reflect*) can only be initiated by the headquarter agent and sent to the human. The reflection consisted of a pre-established message (see Appendix A) intended to increase the human's awareness of the present learning opportunity. It is automatically triggered when the human does not respond to the learning activities by the robot.

4. Testing the Effects of LDPs

Partners of a human-AI teams need to be able to learn from interactions about the task (*mental task model*), about themselves (*self-model*), about the other (*theory-of-mind*), and about their relationship (*interdependence*). In section 2 we discussed the requirements for co-learning to occur, and distinguished between *emergent* learning (learning that evolves naturally from using the results on spontaneously initiated adaptations), and *intentional* learning (learning that is the result of designed activities, deliberately introduced to support the team with acquiring the knowledge and skills to perform adequately). In section 3 we explained how we implemented the requirements and guidelines into a testbed for co-learning. This testbed is intended to study the effects of learning interventions (so called Learning Design Patterns (LDPs)) on intentional learning. In the present section we report on the methods and results of an experimental study into the effects of LDPs on (a) the human's experienced quality of collaboration with the AI-partner; (b) the human's learning (in terms of knowledge, awareness, and understanding); and (c) the team's performance.

Table 1: LDPs. Emergent Learning is represented by learning from behavioral cues only. The LDPs concern activities designed to support learning (intentional learning). LDPs that are initiated by the robot are introduced in serial order. Once the human showed evidence of learning, further LDPs were discontinued.

LDP	Activity to induce learning	Purpose
	Provide behavioral cue	Intends to signal participant that something must be done to improve performance
1-Explain	Provide explanation for behavior	Aiming to make partner aware why particular behavior is needed
2-Assign	Give assignment to conduct a specific action	To force particular behavior from partner
3-Reflect	Instruct to reflect	To consider what causes obstacle and about solutions

4.1 Methods

4.1.1 DESIGN

To investigate the effect of LDPs on co-learning, participants were randomly assigned to one of two conditions. Participants in the 'emergent learning' condition had to rely on observations and experience in order to learn (i.e. 'behavioral cues only'). Participants in the 'intentional learning' condition were provided with LDPs that initiated activities intended to support learning. Figure 5 shows the design, the presentation of materials and measurements.

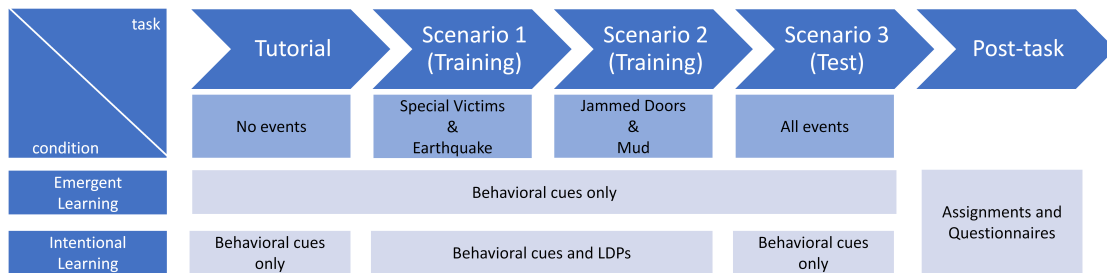


Figure 5: experimental design and procedure

4.1.2 PARTICIPANTS

Participants were recruited from two sources: a participant database and the Netherlands Defense Academy (NLDA). The following inclusion criteria were used: between 18-55 years of age; a college degree or military education (or currently following this education). In total, 45 participants took part (21 male, 23 female, 1 undisclosed). Participants were asked to indicate their age on an age-category scale, with each scale point comprising 5 years (e.g.

20-25; 26-30 and so on). 80% of the participants were younger than 35. Participants were asked to indicate their self-assessed experience with computer games. 75% of the participants rated their experience as low-to moderate. Post-hoc analyses showed that there were no differences between conditions in terms of participants' gender, age and game-experience (see Appendix G for details).

4.1.3 MATERIALS

A simulation of the USAR task (see section 3.1) was developed using the Python programming language (version 3.7) and the MATRX (multi-agent teaming rapid experimentation)² 2.3.0 software library. The MATRX software provides a framework for rapidly building 2D top-down, grid-based environments in which multiple human and autonomous agents can work together collaboratively in a laptop environment. The participant played the role of human explorer in the team; an autonomous agent was programmed to play the role of robot explorer. The robot behaved according to the rules presented in section 3.3. The task was presented using a laptop. Participants controlled the avatar of the human explorer using the four arrow buttons. They could select LDPs (when applicable) by using the mouse. A video was made to instruct participants about the task and experiment. All materials were in Dutch.

4.1.4 PROCEDURE

The participants conducted the experiment individually, in batches of four, under supervision of an experiment leader. Participants sat behind a table with a laptop, external keyboard, HR-monitor, and mouse. Booklets with information and questionnaires were on the table, in closed position. The experiment supervisor sat in the back of the room, with a clear view on all four work stations.

First the experiment leader asked participants to read the information about the experiment, gave answers to any questions asked, and requested participants to sign the informed-consent form. Then the experiment proper started.

Participants performed four runs, as shown in Figure 5: a tutorial by means of a video-clip; then two training sessions with two learning challenges (see section 3.2); a test session including all learning challenges; followed by a post-task session. At the end, the experiment leader gave a debriefing.

4.1.5 MEASURES

Effects of LDPs on co-learning were evaluated with respect to (a) the human's experienced quality of collaboration with the robot; (b) the human's knowledge, awareness, and understanding of its robot-partner; and (c) the team's performance. Table 3 shows all measures used in the experiment, and Appendices C and D report all items of the measures.

Collaboration: The participants' evaluation of the collaboration with the robot was measured by using a Dutch translation of the collaboration fluency questionnaire (Hoffman, 2019). In that questionnaire, participants rate 24 statements about their collaboration with the robot on a 5-point scale. Five dimensions of collaboration fluency are distinguished: *human-robot fluency*; *relative contribution of team members*; *trust in the robot*; *positive teammate traits*;

2. <https://matrx-software.com/>

performance improvement over time; and working alliance.

Team Feeling: Four questions were constructed to measure (a) the participant's monitoring of the robot's behavior; (b) the participant's monitoring of the chat communication with the robot; (c) to what extent the participant felt a teaming relationship with the robot; and (d) how useful the participant considered the collaboration with the robot for the team to complete its task.

Learning: To investigate the participant's understanding of the task, of its own role and that of its partner, and of the team relationships we administered, -for each of the four learning challenges (see Table 1)-, a dedicated combination of questions and assignments. For most of the questions, participants were asked to motivate their answers. Figure 6 provides an example of an assignment/question combination that intends to measure whether the participant understands that the robot uses its experiences to further develop its understanding of the task and of the team. In the initial instruction it was explained to participants that at the start of their collaboration the robot was trained to execute its own tasks for the mission, but that it had not yet experience with performing in a team. Thus before commencing the task, the robot knows what buildings are; what a victim is; it knows that a victim needs to be lifted and carried to the command post, et cetera. However, it has not been trained with human explorers yet. So tasks that require team coordination have to be learned during collaboration. In the example question we request our participant to take the perspective of a not yet fully-trained robot, and then ask how the robot will behave.

Team Performance: Team performance was measured by logging the efficacy (i.e., number of victims saved) and the efficiency (e.g., task duration; idle time; navigational choices) of the partners (see Table 2).

4.1.6 ANALYSIS OF RESPONSES TO OPEN QUESTIONS

Participants' motivations to the open questions on collaboration and learning (see Appendix D) were analysed using Thematic Analysis (Kiger & Varpio, 2020). Two researchers engaged in open and semi-closed coding of the answers, followed by clustering exercises of the given codes. The following steps were taken:

- Both researchers open-coded the same random sample of five participants
- The researchers engaged in a discussion to agree on a coding strategy, deciding the level of detail and focus of the codes
- Both researchers semi-closed coded half of the questions, using the agreed coding strategy
- Both researchers went through the codes given by the other researcher, thereby clustering and regularizing codes
- Any changes made to the initial codes were then thoroughly discussed by the researchers, until disagreements were resolved
- The researchers evaluated the final codes given for each question, sorted by group, and summarized the main themes.



Figure 6: Example of a question eliciting the model that the participant has of the robot (i.e. Theory of Mind).

Imagine that you and the robot just started for the first time with performing as a team. Now look at the image that shows a possible situation that may arise.

- a. Draw the route you think the robot will take from location A to location B.
- b. Why do you think the robot will take this route?

4.2 Results

Effects of LDPs were evaluated with respect to (a) the human’s experienced quality of collaboration with the AI-partner (section 4.2.1); (b) the human’s learning (in terms of knowledge, awareness, and understanding) (section 4.2.2); and (c) the team’s performance (section 4.2.3). Analysis of the question whether our human-robot task yielded a proper environment for the study of co-learning is reported in the Discussion.

4.2.1 EFFECTS ON EXPERIENCED COLLABORATION

It was hypothesized that the participants engaging in explicit learning activities (*intentional learning group*) had more opportunity to develop thorough understanding of the robot partner, resulting in a more fluent collaboration. Results on the collaboration fluency questionnaire (Hoffman, 2019) showed no evidence for this. No differences in mean scores were found on any of the subscales: *human-robot fluency*; *relative contribution of robot to team*; *trust in the robot*; *positive traits of teammate*; *performance improvement over time*; and *working alliance* (respective over-all means and sd’s: 3.28 (.86); 2.93 (.69); 3.40 (.84); 3.71 (.70); 3.52 (.99); 3.45 (.71)).

It was also hypothesized that engaging in learning activities would increase participants’ communication with the robot; more attention to the robot’s actions; an increased feeling of

Table 2: Measures that were used to test effects of LDPs on experienced collaboration, learning, and team performance.

Measure	Unit	Description
To test how participants experienced the collaboration		
Collaboration fluency (Hoffman, 2019)	Likert scales (1-5)	
Monitored chat	Likert scale (1-5) + open	
Monitored robot	Likert scale (1-5) + open	
Felt as team	Likert scale (1-5) + open	
Felt benefit of collaboration	Likert scale (1-5) + open	
Mental taxation	Scale (1-20)	
To test what partners learned		
Robot learned carry	Yes/no	
Robot learned to avoid mud	Yes/no	
Learning results - Mud tiles	Custom test questions	
Learning results - Carry together	Custom test questions	
Learning results - Taking shelter	Custom test questions	
Learning results - Remove broken doors	Custom test questions	
To test team performance		
Task duration	Simulation ticks	time to task completion
Victims saved	Count	number of victims brought to the command post
Idle time	Simulation ticks	Accumulated time of human standing still
Mud moves by human	Count	Number of mud tiles traversed by the human
Mud moves by robot	Count	Amount of mud tiles that were traversed by the robot
Hit by earthquake	Count	Number of times that the human was not taking shelter during an earthquake (and therefore being hit)

being one team with the robot; and an improved appreciation of the robot's contribution to the team. The results did not reveal an increased feeling of being one team, nor of an improved appreciation of the robot's part in the team's mission, as participants of the *intentional learning condition* rated these about the same as participants of the *emergent learning condition* ($t(43)=-.87$, n.s.; and $t(43)=-.85$, n.s.), respectively).

Participants of the *intentional learning group* responded to pay more attention than participants of the *emergent learning groups* to what the robot communicated ($t(43)=-2.53$; $p < .01$). However, they did not monitor the robot’s behavior more closely ($t(43)=-.66$, n.s.).

In general, participants did not evaluate the task as particularly taxing (mean of 8.2 on scale between 0-20). From the analyses on participants’ motivation of their answers to assignment, we found that many participants evaluated the task as fairly simple, because of the moderate pace of task execution. Furthermore, they said to experience little to no time pressure. However, the cognitive load that participants of the *intentional learning group* experienced was higher than for the other group (6,64 versus 9,74; $t(43)=-2.55$, $p < .01$).

4.2.2 EFFECTS ON KNOWLEDGE AND UNDERSTANDING

It was hypothesized that LDPs positively affect participants’ knowledge, awareness, and understanding about the task, themselves, and of their robot partner. Knowledge and understanding was assessed with measures (see Table 2) that correspond to each of the four learning challenges (see Figure 3), namely: (i) moving across mud tiles, (ii) carrying severely wounded victims together, (iii) taking shelter in case of an earthquake, (iv) removing jammed doors.

Navigate to avoid mud It was expected that participants of the *intentional learning group* would be more likely to successfully advise the robot to avoiding mud while navigating than participants of the *emergent learning group*, as the former has explicit interventions (i.e. explanations, assignments) to point out the relationship between mud and navigation speed, whereas the latter group could only do so by demonstrating the appropriate behavior. The number of participants that selected a particular response option for the important mud-related measures are reported in Table 3.

In general, results show that all participants acquired a good understanding of the effect of mud on the human’s navigation speed, as indicated by their responses to the assignment to draw a good route between A and B ("draw_route_A-B"). Almost all participants chose a mud-avoiding route. The participants’ answers to the open questions revealed that they avoided the mud for the right reason: they realized that mud slowed down speed of movement (some of the participants also believed mud to be dangerous). The few participants who decided to navigate through mud were also aware of the slow-down effect, but they said to believe that navigating through mud is nevertheless faster than taking a detour (which was an incorrect belief, as we programmed routes through mud to take longer than mud-evading routes).

A chi-square test for independence showed that 65% of the participants in the intentional learning-group (15 out of 22 participants) noticed the slowing down of the robot significantly more often than participants of the emergent learning-group (32%; 8 out of 23), $\chi^2(1)=5.0$, $p = .025$, Cramer’s $V/\phi = .33$.

We were interested to see whether the participants’ understanding of the robot (i.e. their Theory-of-Mind of the robot) changed during their collaboration. In order to test participants’ knowledge of the robot’s knowledge on the effects of mud on navigation speed, we presented them with a map of the rescue area and asked the participant to draw a route between two buildings that they thought the robot was likely to take. The area between the two buildings was partly covered with mud. See Figure 2 for an example. We asked participants to draw the route they believed the robot would take in the very beginning

Table 3: The number of participants that selected a particular response option for the important mud-related measures (not all participants completed all drawing test questions (for example because their drawing was missing, was unclear, or when participants connected locations other than A and B). This was the case for one participant (intentional learning group) for B1, three participants for B2 (1 from emergent learning; 2 from intentional learning), and one participant for B3 (intentional learning group). These data were not included in the analysis).

Variable	Assignment/questions	Group		
			Through mud	Avoiding mud
Draw_route_A_B	draw route you would take to go from location A to B	Emergent	3	19
		Intentional	1	21
			Not noticed	Noticed
Noticed_mud_slowing_down_robot	Did you notice that robot moved slower when navigatingthrough mud?	Emergent	15	7
		Intentional	8	15
			Through mud	Avoiding mud
Draw_robot's_route_at_beginning	draw route you think that robot would take from A to B at beginning of collaboration	Emergent	8	13
		Intentional	13	8
			Through mud	Avoiding mud
Draw_robot's_route_at_end	draw route you think that robot would take from A to B at end of collaboration	Emergent	8	14
		Intentional	6	16

of their collaboration ("Draw_robot's_route_at_beginning"). Then we presented another map with a different configuration of buildings and mud, and asked participants again to

draw the route they believed the robot would take at the end stages of their collaboration ("Draw_robot's_route_at_end").

Results show that 59% of the participants of the intentional learning condition (13 out of 21) understood that the robot would cross through mud in the beginning of their collaboration (when the robot had not yet learned the impact of mud on speed of movement), whereas in the emergent learning group only 36% of the participants seem to be aware of the robot's learning during the collaboration (8 out of 21). Understanding the implications of the initial knowledge gap of the robot suggests a Theory-of-Mind capacity of the participants. A chi-square test for independence showed the differences between groups not to be significant, $\chi^2(1)=2.38$, $p = .108$, Cramer's $V/\phi = .238$. The qualitative analysis of participants' responses to their motivation of the route they predicted revealed that various participants of the intentional learning group explicitly mentioned that the robot learned to avoid mud over time, whereas participants of the emergent learning group did not express this explicitly, and indicated not to be sure as to whether and how the robot changed its behavior on this aspect over time.

Results on predictions of the robot's route at the end of their collaboration showed no between groups difference. Still 32% of the participants (14 out of 44) predicted the robot to take a route through mud.

Carrying severely wounded victims together It was expected that teams of the intentional learning group would be more likely to learn that severely wounded victims need to be carried jointly by human and robot. The number of participants that selected a particular response option for the important carry-victims-together measures are reported in Table 4.

When asked whether they were aware that the robot initially did not know that severely wounded victims should be jointly carried, eight participants in the *intentional learning group* indicated to be aware of this, and nine of the *emergent learning group*. Hence no differences between groups were found on this measure.

When the participant stood still next to a victim waiting for help, this was regarded as giving a cue to the robot to provide assistance (the participant may have been aware of unaware of giving this cue). Almost all of the participants demonstrated this cueing behavior required to make the robot learn (100% for the *intentional learning group*; 96% of the *emergent learning group*). In addition to providing behavioral cues, participants could also give cues to the robot by means of the communication chat. Results show that 52% of the participants of the *intentional learning group* (12 out of 22) expressed to send messages to the robot, whereas only 23% (5 out of 22) of the participants of the *emergent learning group* did. A chi-square test showed the difference to be significant, $\chi^2(1)=4.15$, $p = .041$, Cramer's $V/\phi = .304$.

Finally, the content of the communication with the robot was investigated. Participants were asked afterwards to recall their messages to the robot about carrying victims together. About half of the participants in the *intentional learning group* reported having sent a message to the robot about carrying victims together: 47% (8 out of 17) gave an 'explanation', the rest (18%; 3 out of 17) an 'assignment'. Four participants who opted for offering an explanation reported doing so because then the robot can preserve this knowledge and learn from it, the other four said they simply chose the top option of the menu; two of the three participants opting for the assignment reported doing so because it sounded more friendly to them. Of

Table 4: The number of participants that selected a particular response option for the carry-victims-together measures.

Variable	Question	Group	Response		
			Not noticed	Noticed	
Noticed_robot's_knowledge_gap_about_carrying	Did you notice that robot initially did not know that severely wounded victims need to be jointly carried?	Emergent	8	14	
		Intentional	9	14	
			No	Yes	
Sent_chat_message_to_robot_about_carrying_victims	Did you send message(s) to robot on carrying severely wounded vicioms together?	Emergent	17	5	
		Intentional	11	12	
			Explan- ation	Order	Do not remem- ber
Type_of_help_provided	Did you explain to the robot why joint carrying is necessary, or did you order the robot to help with carrying the victim?	Emergent	0	1	4
		Intentional	3	8	1

the *emergent learning group*, very few participants reported to have sent messages. Some responded to assume that the robot can learn regardless of the human's actions; others added to send messages 'out of friendliness towards the robot'.

Taking shelter in case of an earthquake It was expected that participants of the *intentional learning group* would be more likely to learn that the robot has means to predict upcoming earthquakes and to learn how to take shelter when this happens. The number of participants that selected a particular response option for the important measures of the earthquake-learning-challenge are reported in Table 5.

Table 5: The number of participants that selected a particular response option for the important taking-shelter-for-earthquake measures.

Variable	question	Group	Response	
			Not noticed	Noticed
Noticed_red_cross	Did you notice the red cross appearing on your avatar?	Emergent	0	22
		Intentional	3	20
			Incorrect	Correct
Knowing_when_red_cross_appeared	In what situations did the red cross appear?	Emergent	19	3
		Intentional	9	11
			Incorrect	Correct
Knowledge_of_effects_following_red_cross	What happened after a red cross had appeared?	Emergent	10	12
		Intentional	7	13
			Incorrect	Correct
Knowledge_of_when_an_earthquake_is_imminent	Was it possible to detect an imminent earthquake, and if so, how?	Emergent	14	8
		Intentional	5	18
			Incorrect	Correct
Knowledge_of_robottaking_shelter	draw a cross on the map where you think the robot will be when an earthquake is coming; motivate your answer	Emergent	10	12

Continued on next page

Table 5: The number of participants that selected a particular response option for the important taking-shelter-for-earthquake measures. (Continued)

		Intentional	0	23
			Not in door frame	In door frame
Knowledge_of _whereto_takes _shelter	draw a cross on the map where you would stand when an earthquake is coming	Emergent	13	9
		Intentional	3	20

Results showed that almost all participants noticed the red cross, with no differences between groups. Participants were asked to describe the situation in which a red cross would appear. Results show that 55% (11 out of 20) of the participants of the *intentional learning group* understood the meaning of the red cross, against 15% (3 out of 22) of the participants of the *emergent learning group*. This is a significant difference $\chi^2(1)=8.07$, $p = .005$, Cramer's $V/\phi = .438$. The results of the qualitative analysis on the assignments and open questions show that although most participants understood that the red cross has something to do with the earthquake, the *emergent learning group* generally reported that the red cross roughly appears around the moment of the earthquake. Participants of the *intentional learning group* were able to indicate this much more precisely, explaining that the appearance of the red cross was also dependent upon whether or not the human has taken shelter in time to seek protection against the earthquake.

Furthermore, 65% (13 out of 20) of the participants of the *intentional learning group* reported correctly the effects that followed the appearance of a red cross, against 55% (12 out of 22) of the *emergent learning group*, a non-significant difference. A chi-square test showed that participants of the intentional learning group learned better than participants of the other group to predict when a new earth quake was due, 78% vs. 36%, $\chi^2(1)=8.09$, $p = .005$, Cramer's $V/\phi = .424$. In response to open questions, the participants of the *intentional learning group* indicated to be observant when the robot stood still in a door opening, and said they learned from the robot how to seek shelter (by copying the robot's behavior). Participants of the *emergent learning group* were close to finding this out, but they were not fully accurate. Eleven of them pointed to the vibrating screen as being their alarm signal. However, the vibrating buildings are the expression of the earthquake, not its announcing signal. Furthermore, many participants of this group misinterpreted the behaviour of the robot erroneously, e.g., interpreting 'standing still in a door frame' as 'doing nothing', rather than as 'seeking shelter'.

Participants of the *intentional learning group* also learned better than participants of the *emergent learning group* to understand where the robot was standing (100% vs. 60%), $\chi^2(1)=13.44$, $p = .001$, Cramer's $V/\phi = .547$, and learned better to find an appropriate spot themselves to take shelter for an imminent earthquake, (87% vs. 38%) $\chi^2(1)=10.4$, $p = .001$, Cramer's $V/\phi = .481$. These findings are corroborated by the observations that participants of the *intentional learning group* were less often hit by an earthquake (mean=8.5;

sd=3.7) than participants of the *emergent learning group* (mean=11.4; sd=1.7). This is a significant difference: ($t(38)=-6.26, p < .01$). Most participants, as inferred from responses to the open questions, understood that they had to seek shelter against earthquakes in a door frame. Participants of the emergent learning group argued more often to seeking shelter in the headquarters building (which is an unnecessary limitation for shelter-seeking), whereas participants of the intentional learning group more often correctly reported that any door opening would do. This indicates a deeper understanding of the relationship between earthquake, behavior, and consequences. A similar sign of deeper understanding was that participants of the *intentional learning group* mentioned a difference, in terms of safety, between seeking shelter inside a building, and seeking shelter in the door opening. This supports the finding that the *intentional learning group* understood better what truly functions as shelter against earthquakes.

Removing jammed doors It was expected that participants of the *intentional learning group* would be better at learning to jointly open jammed doors. Results regarding this learning challenge were evaluated by means of participants' responses to an assignment and corresponding open questions. The participant was presented with a picture of a damaged building (doors of damaged buildings were always jammed) with a severely wounded victim inside. The participant was asked to describe in detail, step-by-step, what would be needed for the team to evacuate the victim (see Appendix D). Results revealed that almost all participants had learned what to do. The level of detail in their responses varied from participant to participant. Most people explicitly mentioned to pro-actively offer assistance to the robot by helping to remove the jammed door. A few participants indicated to experience miscommunication and bad timing in their joint work with the robot. Overall, all participants managed to learn the need for joint action fairly quickly and accurately. The learning challenge that we aimed to design into the experiment proved to be fairly easy to solve for the participants. Hence, this challenge provided insufficient opportunity to investigate differences between the two learning conditions.

4.2.3 EFFECTS ON TEAM PERFORMANCE

It was hypothesized that the *intentional learning group* would show better team performance than the *emergent learning group*. Participants of the *intentional learning group* appeared to perform better than participants of the *emergent learning group* on some of the measures: fewer moves through mud by the human rescue worker ($M=18.2$ (sd=14.4) vs. $M=24$ (sd=16.3)); fewer moves through mud by the robot ($M=1.2$ (sd=3.1) vs. $M=4$ (sd=11.3)); and fewer hits to the human rescuer by earthquakes ($M=8.5$ (sd=3.7) vs. $M=11.1$ (sd=1.7)) (i.e., did not seek shelter during an earthquake). However, independent-samples t-tests on the measures revealed that the *hit_by_earthquake* measure was the only significant difference ($t(30.9) = 3.48, p < .05$).

The measure *victims saved* was not analyzed, because all participants who completed the mission must, by default, have saved the same number of victims. As all our participants completed the mission in time, this measure was dropped from analyses.

5. Discussion

An important topic in human-machine teaming is the question how humans and AI-based agents can develop and enhance their knowledge about their joint task and about each other (Demir, McNeese, & Cooke, 2020). For a hybrid team to function effectively, both human and AI must develop and maintain internal representations that encompass knowledge about the task context, their own responsibilities, and the roles of others. Furthermore, partners must be able to apply this knowledge in a variety of task contexts. The process of teammates collaborating, adapting to each other, and understanding each other is called co-learning (Van Zoelen et al., 2021a). This co-learning may occur implicitly and naturally, emerging from partners' experiences on task. Co-learning may also take place intentionally, in purposely designed situations (Schoonderwoerd et al., 2022).

This paper presents an analyses of the requirements for co-learning to occur. These requirements pertain to the design of a task and the task environment; the functionalities and behavior of the AI-based agents; and to the qualities of the human team member. This is discussed in the next section. The paper also reports an empirical study into the effects of deliberately designed learning interventions on learning and performance. This is discussed in section 5.2.

5.1 Requirements for Human-AI Co-Learning

When considering co-learning in human-AI teams, it should be taken into account that human and AI team members have different kinds of mental models, embodiments, and ways of learning. An environment that intends to foster mutual learning should accommodate for these fundamental differences. There is a need to explore the conditions under which co-learning can flourish.

It has been argued that dependency is such a critical condition (Johnson et al., 2014; Johnson & Bradshaw, 2021; Van Den Bosch et al., 2019). Dependency compels team members to collaborate, to coordinate activities, and to provide or request help. Engaging in these activities enable partners to develop and refine their internal representations. This all is co-learning: a process in which collaborating partners adapt to each other and learn together over time (Van Den Bosch et al., 2019).

In this study we adopted the notion of interdependency as the principal construct for developing an experimental testbed for empirical research into co-learning. We added other important notions, such as the principle that learning should take place in the context of the collaborative task itself, rather than occurring separately (outside the task). Furthermore, as learning is a dynamic and continuous process, we chose to provide multiple opportunities for achieving a learning goal (rather than aiming for single-shot learning). Another important precondition for co-learning concerns the mind-set of the partners, in our study the mind-set of the human partner. It is important to design the environment in such a way that the human considers the interactions with the AI-based partner as plausible (Cross & Ramsey, 2021). To create the conditions for co-learning in our testbed, the human should regard the AI-based robot as a true partner, should be prepared to improve as a team by joint collaboration, and should have the willingness to learn about the team partner. Finally, we also added additional activities to the testbed, designed to enhance learning, such as providing opportunities for partners to explain their needs or request to the partner (explanations), to assign requests that will lead to successful collaboration (assign), and to issue calls for

reflection that foster perspective taking and mental model refinement (reflection). Armed with the above requirements, we developed the testbed (see section 3) suitable for empirical research into co-learning (see section 4).

Our measures and observations of participants during their collaborations with the AI-based robot allow us to draw the following conclusions:

With respect to interdependency, the hard dependencies intentionally designed into the testbed proved to stimulate participants to collaborate with the robot. In response to questions administered after the experiment, participants acknowledged the need for collaboration as enforced by the hard dependencies. With respect to soft interdependencies, participants required ample opportunities to understand the need of collaboration, but progressively started to behave in a pro-active fashion, e.g., by taking position in a door frame, even before the robot did so.

Our testbed enabled all participants to achieve the four learning objectives, although not by all participants at the same pace and not by using the same arguments in their motivations. The learning challenges could be solved by direct interactions between humans and robots, perhaps making learning somewhat too easy for research purposes. In future research we intend to include more difficult learning challenges that require making inferences based upon incomplete and/or uncertain information. For example, the human observes that two or more AI-based systems demonstrate unexpected behavior, and the participant has to infer what agreements may have been made by the AI-based systems that account for the way they behave. This would require the human to construct a theory-of-mind of the AI-systems from incomplete information.

The responses to questions after the experiment showed a recurrent theme, namely that participants were constantly generating hypotheses to make sense of the behavior of the robot. For example, one participant observed the robot moving through mud, and subsequently made the assumption that the robot is allowed to go through mud, whereas the human is not. Another participant assumed that the robot doesn't need to be informed about the slowing-down effect of mud, as it will keep away from mud spontaneously. This tendency to constantly make assumptions from observations may be explained by the need that people feel to achieve coherence in their mental model of the task world. When the behavior of the robot doesn't make sense to them, people restore sense by generating assumptions. This also occurs in all-human teams (Driskell, Salas, & Driskell, 2018), although unexpected behavior is perhaps more likely to occur in human-AI teams. Therefore, in human-AI teaming tasks, it is important that humans are taught to think critically about any unexpected behavior of their partner(s), and to critically test any assumptions they make. At the same time, it demonstrates the importance of making AI-based systems transparent (e.g., by explaining themselves), of making explicit and clarified requests, and to initiate shared after action reviews by the team to facilitate learning (Anjomshoe, Najjar, Calvaresi, & Främling, 2019).

An important condition for accomplishing co-learning is that the human perceives the AI as a real team partner, with a human-like mental model and a sense of self. If not, humans may regard the AI-based system as untrustworthy and reject it as a team member (Groom & Nass, 2007). Furthermore, when humans feel a shared commitment with the AI-based partner, this improves the intention to collaborate (Correia, Mascarenhas, Prada, Melo, & Paiva, 2018; Correia, Melo, & Paiva, 2024). Our measures and observations reveal that participants demonstrated more interaction with the robot, and felt more shared commitment

than in our previous studies. This time, participants considered the contribution of the robot to the team’s performance as crucial. They also paid much more attention to what the robot was doing than in previous efforts. We attribute this to our efforts to design multiple human-robot interactions in the task, and to ensure symmetry in the communication (e.g., robots always replied to messages from the human; and robots also initiated communication with the human). This improved commitment and collaboration is a step in the right direction, although more efforts are needed to establish a true and collaborative team spirit between humans and AI-based systems.

5.2 Effects of Learning Interventions

We developed Learning Design Patterns (LDPs) to organize activities and interactions within a human-robot team, with the objective to support partners’ learning about their joint tasks and about each other. We constructed measures (e.g., specific and open questions, questionnaires, assignments, performance logs) to investigate the effects of LDPs on participants’ collaboration experiences, knowledge and understanding, and team performance.

5.2.1 COLLABORATION

Although responses to open questions suggest that participants of the intentional learning group were more responsive and open to their robot partner, no between-groups differences were found on any of the subscales of the collaboration fluency scale (Hoffman, 2019). Thus, the learning activities seem to have no effect on how collaboration with the robot is experienced. However, it may also be that the collaboration fluency scale, administered at the end of the experiment, may not be sensitive enough to reveal differences in teamwork fluency as experienced earlier in the experiment. Moreover, most participants filled out the questionnaire very quickly, and they very rarely used the extreme values of the scales. To evaluate how humans experience the collaboration with an AI-based partner it is therefore worthwhile to consider alternative methods, like measures obtained during the collaboration itself (rather than administering a questionnaire afterwards). Measures may also be more specifically targeted at the collaboration elements during the task, rather than requesting to evaluate generally formulated statements, as in the CF-questionnaire.

The finding that participants of the *intentional learning group* experienced a higher cognitive load may be due to the demands of engaging in the learning activities.

5.2.2 KNOWLEDGE AND UNDERSTANDING

We designed four learning challenges with associated learning objectives to study the effects of LDPs on the participants’ knowledge and understanding. One of the challenges (learning to provide assistance with opening jammed doors) proved to be too easy, which made studying the effects of LDPs impossible. Results show that the LDPs for the other challenges produced desired outcomes, in particular: taking more notice of the robot’s actions (i.e., better noticing of slowing down of robot on mud-covered areas); maintaining more frequent communication with the robot (in all three challenges); a better understanding of the task and implications of actions (i.e., why, when and where to take shelter in the earth quake learning challenge); and occasionally also better task performance (i.e., finding appropriate shelter against earthquakes more often).

5.2.3 PERFORMANCE

The learning induced by the learning activities appeared to also bring about a better team performance, as participants in the *intentional learning group* had better mean values on several performance measures (i.e., smarter navigation, fewer hits by earthquakes). However, differences with the *emergent learning group* were significant for only one measure. Thus, results provide no conclusive evidence that improved knowledge and understanding lead to better performance. However, the finding that participants of the *intentional learning group* experienced a higher cognitive load indicates that better performance may be achieved if training would be sustained over a longer period of time. It is known that learning requires mental efforts that do not necessarily transfer immediately into better performance (Schmidt & Bjork, 1992). So it may well be that the higher mental load reported by participants of the *intentional learning group* was caused by learning processes going on, but that the exposure to training was insufficient to attain a higher performance level.

6. Conclusion and Future Work

The rapid growth of artificial intelligence will bring forth new ecosystems in which humans and AI-driven technology act as complementing partners in hybrid teams (Peeters et al., 2021). For this to be successful, the conditions must be created in which partners jointly learn to recognize, acknowledge and utilize their respective capabilities (Van Den Bosch et al., 2019). In this paper we addressed the question how to invoke, research, and instigate co-learning in a human-AI team. Our analyses of how to design environments for human-AI co-learning reveal that considerable efforts are needed to ensure mutual understanding, taking into account that humans and AI-based partners have different kinds of world representations, embodiments, languages, and ways of learning. Co-learning environments should accommodate partners to exchange information that is relevant to the task and to the team. Such information exchange may be explicit (e.g., verbally) and implicitly (e.g., behavioral cues), and will support shared understanding. In addition, we identified the need for partners of a human-AI team to achieve a shared commitment, in which the human perceives the AI as a true team partner. This aspect of cohesion is very important in human teams (Salas, Shuffler, Thayer, Bedwell, & Lazzara, 2015), and also for human-AI teams (Beans, 2018; Correia et al., 2018, 2024). In our present study we encouraged shared commitment by ensuring symmetry and complementarity in the behavior of partners, and by designing incentives in the task for cooperation and communication between partners. In comparison to our previous learning environments (Schoonderwoerd et al., 2022) this resulted in more interactive human-robot teamwork. Human participants used their knowledge of the robot’s capabilities and limitations by demonstrating pro-active behavior, and reported to experience being a team with the robot. This is a promising finding. However, the nature of the human-robot collaboration that we observed in our study is still far from what we observe in newly formed human teams. Further research is needed to develop a framework for human-AI teams where: (a) the value of learning about each other is intrinsically embedded in the context of the task; (b) communication between teammates is facilitated so that they can share directives, explanations, and reflections pertaining to their behavior; (c) learning interactions occur dynamically, back and forth over a longer period of time; and (d) partners have sufficient opportunities to develop, refine, and validate their internal representations.

We found promising outcomes of the interventions designed to initiate and improve learning in a human-AI team, in particular on human's knowledge development and understanding. Our study has also made clear that such learning does not automatically and rapidly lead to better task and team performance. Longitudinal research that provides prolonged training in collaboration is needed to shed more light on how humans and AI-based partners grow over time from adapting to each other, to understanding each other, to performing as an expert team (Van Zoelen et al., 2021b).

In our next study into human-AI co-learning we intend to foster and utilize metacognitive skills to support the human and the AI-partner with figuring out what processes underlie events in the team's task and the behaviors of partners. This typically includes: recognizing patterns in collaborating interactions, reasoning, developing and verifying hypotheses, taking different perspectives, and other skills that contribute to a better understanding. By sharing and validating their accumulated understanding, humans and AI-based partners will better learn how to work together as a team.

Conflict of Interest

On behalf of all authors, the corresponding author states that there is no conflict of interest.

Appendix A. Pre-defined Communication Messages

Messages that the robot could send to the human participant (translated from Dutch):

- "I will start at the Post Office, can you start at Dorpsstraat 1?"
- "I am coming to help you."
- "Okay"
- "I cannot pass through the door, I cannot open it myself."
- "Can you help me open the door?"
- "I am taking shelter for an upcoming earthquake."
- "Can you walk to the nearest doorway?"

Messages that the human participant could send to the robot (translated from Dutch):

- "I am coming to help you now."
- "Can you help me carry this victim?"
- "Heavily wounded victims need to be carried together!"
- "Can you please avoid mud?"
- "Mud slows you down."

Messages that headquarters could send to the human (translated from Dutch):

- "We were watching the mission, and wondered: why does the robot not enter this room?"
- "We see that the robot is not coming to help you. Why would that be?"

- "We see that the mission is not progressing very quickly. Why is the robot walking slowly?"
- "We see that the robot is standing still in the doorway. Why is it doing that?"

Appendix B. Text of Instruction Video

[Translated from Dutch]

Explanation of first area

A violent earthquake recently occurred in the region you are in. Fortunately, the buildings and roads are still recognizable, but there will probably have been many victims. The crisis team headquarters set up a rescue operation as quickly as possible.

Various rescue teams work together to bring all victims in the area to safety as quickly as possible. Together with a robot you are part of team *Blauw*. The robot is equipped with smart technology that allows it to autonomously search for victims in a disaster area. This is the first time you and the robot work together. There are four areas assigned to your team. Once you have brought all the victims from an area to the home base, you move to the next area, which is a little closer to the epicenter of the earthquake. You and the Robot must discover the best way to work together in each area. This is possible if you pay close attention to each other and coordinate how you work.

Your team's mission is to search all buildings in the area together and bring found victims to their home base. It is important that you work together well so that earthquake victims are found as quickly as possible and brought to safety.

In this video you see the first area that you and the robot will search. You are the rescuer with the yellow helmet. You can move around the environment with the arrow keys. This is the robot. The robot moves independently in the area.

All these red buildings were affected by the earthquake and need to be searched. The green building, the command post, is your home base, found victims must be brought here. On the right side you will see a chat window, try to keep an eye on this during the mission, the robot may ask you something or messages may appear from headquarters. As soon as the scenario starts, you can check whether victims are in a building by standing in the doorway of the building. Found victims are visible to the entire team and remain visible in the area. So as soon as you have found a victim, this is also known to the robot, and vice versa. Victims can be slightly wounded, orange, or seriously wounded, red. It is also possible that a victim has already died, in which case it will be gray in color. Deceased victims do not have to be taken to the command post.

You can pick up a victim by standing right next to it and right-clicking on the victim. Select the option "Carry slightly wounded victim" from the menu. The victim is then in your arms. Use the arrow keys to walk to the command post with the victim in your arms. You can put the victim down by right-clicking on yourself and choosing the option "put down the victim" from the drop-down menu. Your teammate, the robot, is also able to find victims, pick them up and take them to the command post. The robot does this automatically. After this video the first rescue operation starts. Together with the robot, try to bring all victims to safety as quickly as possible.

Second area (training scenario 1)

Now that you have searched the first area and brought the found victims to safety, it is time to move on to the next area, which is closer to the epicenter of the earthquake. There are immediate differences with the previous area. There are also seriously wounded victims here and the chance of an aftershock also increases. Seriously wounded victims, identified by their red color, must be transported extra carefully on a stretcher. These victims cannot be carried alone. You and the robot will therefore have to work together and lift the victim together. If you carry a victim together, you control the stretcher with the victim using the arrow keys. You can put the victim down again by clicking on yourself with the right mouse button and choosing "put the victim down together". The robot then automatically continues its work. As in the previous area, slightly wounded victims, colored orange, can be carried alone and deceased victims, colored gray, do not have to be taken to the command post.

In the command post, victims are treated for their injuries. As long as victims are in the other buildings, their condition deteriorates. That's why you get to work quickly, because the sooner the victims are in the command post, the greater the chance of survival.

Third area (training scenario 2)

You will move on to the next area on your schedule. The consequences of the earthquake are even greater here. There are more damaged buildings and the damage appears to be more serious. The earthquake caused mud to flow from the slopes into the area, recognizable by the dark gray areas. Some buildings have also collapsed, which can be seen by the broken walls of a building. Headquarters has banned rescue workers from entering collapsed buildings. The robot is allowed to enter these buildings. He can search the collapsed buildings and bring out wounded victims. He can only drag the seriously wounded victims to the door, from there they must be carried further together, just like in the previous scenario. As in the previous area, slightly wounded victims can be carried alone. Due to the conditions in this area, an aftershock is unlikely. You will quickly get back to work to bring all victims to safety.

Fourth area (test scenario)

This last area on your list is in the worst shape. You are near the epicenter of the earthquake and Headquarters expects that there will be many casualties here. There is a high risk of aftershocks and many buildings have collapsed. Try to apply all the knowledge you have learned from previous collaborations with the robot. The robot has also learned from the previous collaboration and will bring its experiences to this area. You will get started quickly.

Appendix C. Scales and Items used of Collaboration Fluency Questionnaire

[Original items of Collaboration Fluency Questionnaire (Hoffman, 2019). Dutch translations of these items were used in the experiment]

Human-Robot Fluency

- The human-robot team worked fluently together.
- The human-robot team's fluency improved over time.
- The robot contributed to the fluency of the interaction.

Robot's relative contribution

- I had to carry the weight to make the human-robot team better. (R)
- The robot contributed equally to the team performance.
- I was the most important team member on the team. (R)
- The robot was the most important team member on the team.

Trust in robot

- I trusted the robot to do the right thing at the right time.
- The robot was trustworthy.

Positive Teammate Traits

- The robot was intelligent.
- The robot was trustworthy.
- The robot committed itself to the task

Improvement

- The human-robot team improved over time.
- The human-robot team's fluency improved over time.
- The robot's performance improved over time.

Working Alliance for H-R Teams

- The robot and I understand each other.
- I am confident in the robot's ability to help me.
- The robot and I trust each other.
- The robot perceives accurately what my goals are.
- The robot does not understand what I am trying to accomplish. (R)
- The robot and I are working towards mutually agreed upon goals.
- I find what I am doing confusing. (R)

R = reverse scale item

C.1 Additional Questions

In addition to the Collaboration Fluency Questionnaire (Hoffman, 2019), four additional questions were administered measuring (a) the participant's monitoring of the robot's behavior; (b) the participant's monitoring of the chat communication with the robot; (c) to what extent the participant felt a teaming relationship with the robot; and (d) how useful the participant considered the collaboration with the robot for the team to complete its task.

[Translated from Dutch]

- During the task I always kept an eye on the robot.
- During the task I always kept an eye on the chat messages.
- I felt like the robot and I were a team together.
- I felt that it was useful for the team's task to collaborate with the robot

The following question was administered to measure how demanding the participant experienced the task:

- Please indicate how mentally taxing you found the task.
(take into account the efforts needed to set priorities, to carry out the actions, to communicate with the robot, and other activities that you performed.)

[Participants were requested to rate this question on a 21-points scale, ranging from 'low mental load' to 'high mental strain']

Participants were requested to explain their answer to each question.

Appendix D. Questions to Assess Learning

D.1 Navigating in muddy area

Question 1

Take a good view of the situation depicted in Figure 7. Draw the route *you* would take from A to B.



Figure 7: Draw the route *you* would take from A to B.

Please explain why you would choose this route.

Question 2

Imagine that you have never worked with the robot on the task before. This is the first time you work together with the robot.

Now take a good view of the situation depicted in Figure 8. It shows an example of a situation that could arise. Draw the route that you think *the robot will take* from location A to location B.



Figure 8: Draw the route that you think the robot will take from location A to location B.

Please explain why you think the robot would be taking this route.

Question 3

Now please imagine that you have worked with the robot on this task before.

Take a good view of the situation depicted in Figure 9. It shows an example of a situation that could arise. Draw the route that you think *the robot will take* from location A to location B.

Please explain why you think the robot would be taking this route.

Question 4

The robot moved slower as it went through the mud. Did you notice this while performing the task? (Please answer honestly, there is no wrong answer).

- Yes, I noticed
- No, I didn't notice

Question 5

Did you say or ask the robot anything about the mud paths during the task?

- Yes
- No

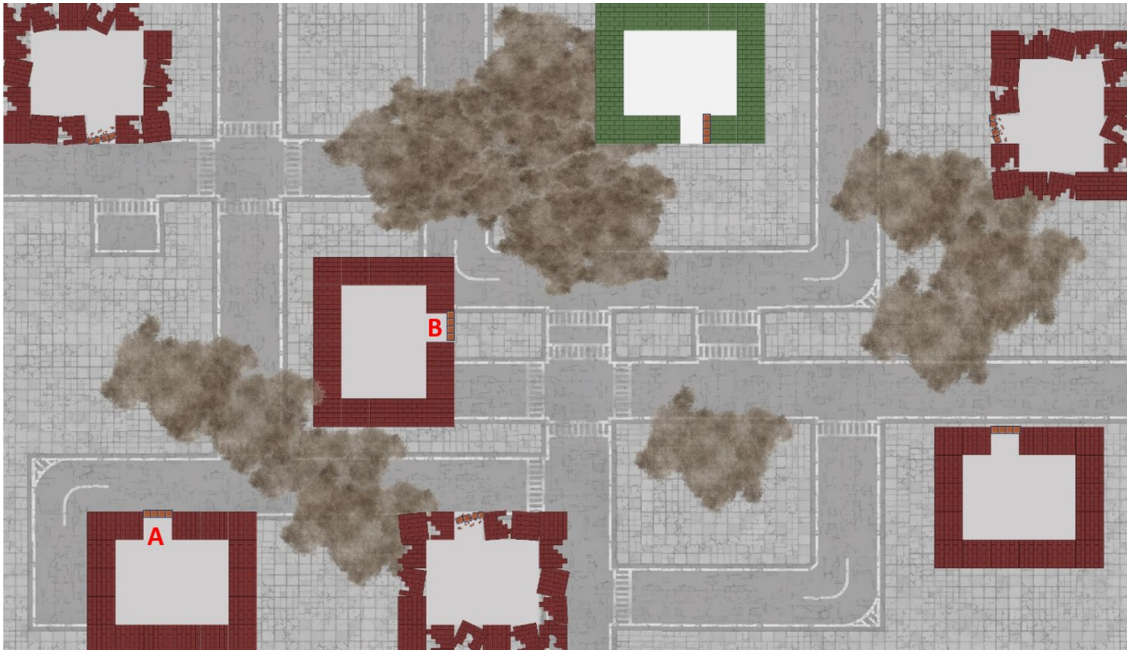


Figure 9: Draw the route that you think the robot will take from location A to location B.

Question 6

What prompted you to communicate with the robot about muddy paths? (please answer truthfully; there is no wrong answer)

- I noticed the option when I clicked on the robot.
- I noticed that the robot moved slower when there was mud, so I tried whether I could ask or say something to the robot by clicking on it.

Question 7

Take a good view of the situation depicted in Figure 10. Which of the two messages did you send first to the robot?

- I asked the robot to avoid muddy paths.
- I told the robot that mud slows down its speed.
- I can't remember

Please explain why you chose this message.

Question 8

Do you think a message about how to navigate when there is mud will cause the robot to exhibit appropriate behavior in the future?

Please explain your answer.

D.2 Dislodging a collapsed door

Question 9

Take a good view of the situation depicted in Figure 11. The victim in building A needs to

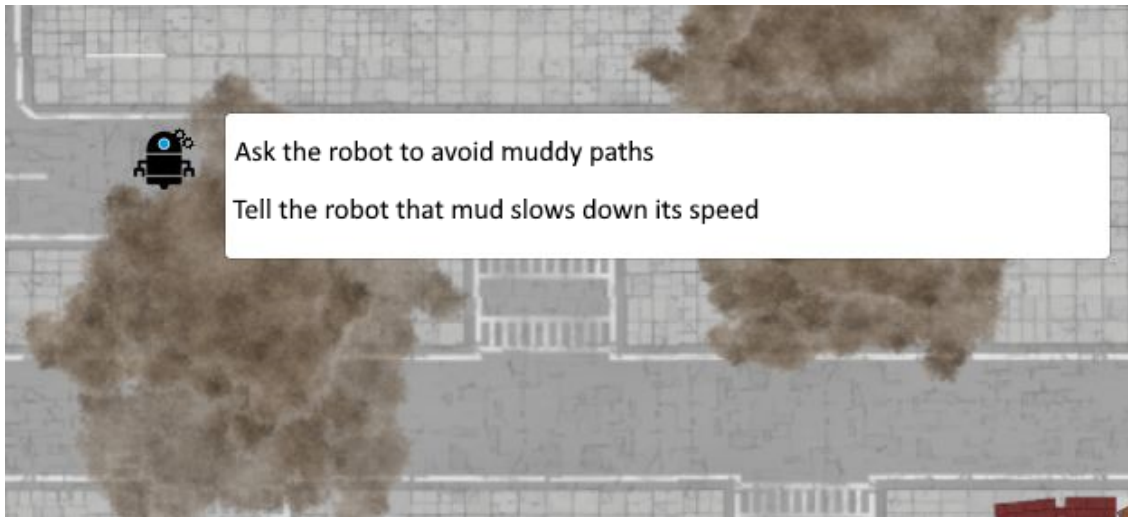


Figure 10: Which of the two messages did you send first to the robot?



Figure 11: The victim in building A needs to be rescued and the robot is almost arriving at this building.

be rescued and the robot is almost arriving at this building.

Describe step by step how your team brings this victim to home base. Try to include all actions in the description.

Question 10

Take a good view of the situation depicted in Figure 12. The image shows a building with a

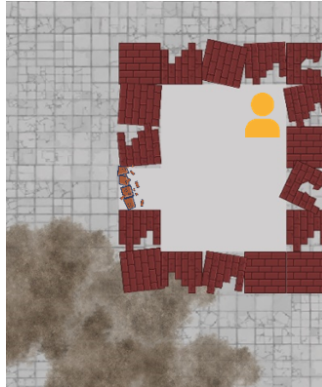


Figure 12: The victim in building A needs to be rescued and the robot is almost arriving at this building.

victim inside. Describe step by step how your team brings this victim to home base. Try to include all actions in the description.

D.3 Carrying a severely wounded victim

Question 11

Take a good view of the situation depicted in Figure 13. Describe step by step how your team

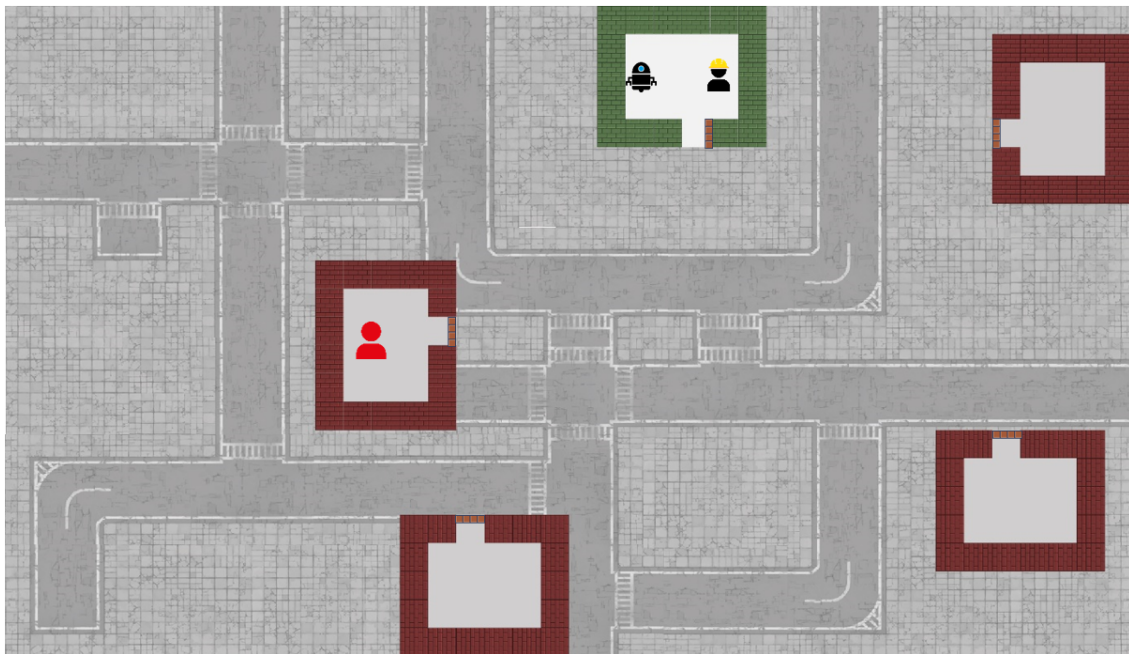


Figure 13: Describe step by step how your team brings the victim in the middle-left building to home base

brings this victim to home base. Try to include all events and actions in the description.

Question 12

Take a good view of the situation depicted in Figure 14. Describe step by step how your team



Figure 14: Describe step by step how your team brings the victim in the lower-left building to home base

brings this victim to home base. Try to include all events and actions in the description.

Question 13

A seriously wounded victim could only be carried in collaboration with the robot. Did you notice that the robot wasn't aware that both of you were needed to carry seriously wounded victims? (Please answer honestly, there is no wrong answer).

- Yes, I noticed
- No, I didn't notice

Question 14

During the task, did you say or ask anything to the robot about lifting a seriously wounded victim?

- Yes
- No

Question 15

What prompted you to communicate with the robot about lifting a seriously wounded victim? (please answer honestly; there is no wrong answer)

- I noticed the option when clicking on a seriously wounded victim

- I noticed the robot to be unaware of the need of carrying seriously wounded victims together. I tried clicking on the victim to see whether I could ask or say something about this to the robot.

Question 16

Take a good view of the situation depicted in Figure 15. Which of the two messages did you

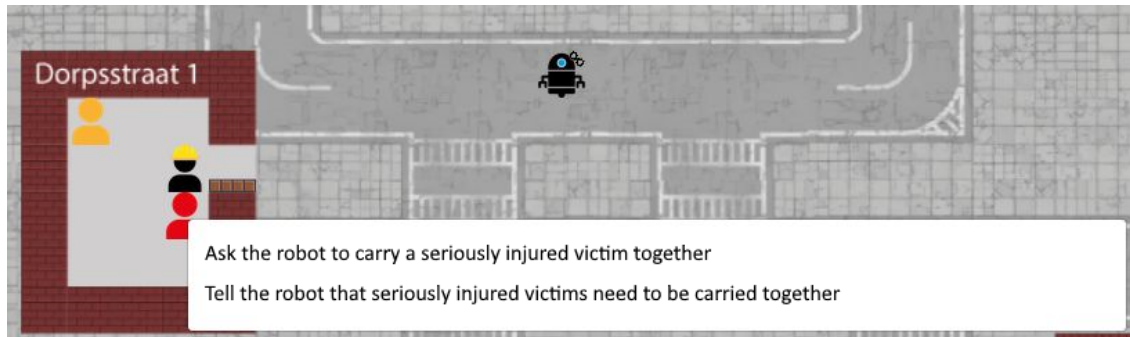


Figure 15: Which of the two messages did you send first to the robot?

send first to the robot?

- I asked the robot to carry a seriously wounded victim together
- I told the robot that seriously wounded victims need to be carried together
- I can't remember

Please explain your answer.

Question 17

Do you think a message about how to carry seriously wounded victims will cause the robot to exhibit appropriate behavior in the future?

Please explain your answer.

D.4 Taking shelter for an earthquake

Question 18

An aftershock occurred at various times during the scenarios. A red cross may have appeared on your avatar (see Figure 16). Did you notice this red cross appearing on your icon?



Figure 16: Did you notice this red cross appearing on your icon?

- Yes, I noticed

- No, I didn't notice

Question 19

Please explain what you think that this red-cross icon means.

Question 20

Explain under what circumstances this red-cross icon appeared.

Question 21

Explain what changed in the environment when this red-cross icon appeared on your avatar.

Question 22

Take a good view of the situation depicted in Figure 17. Mark in the Figure where the robot

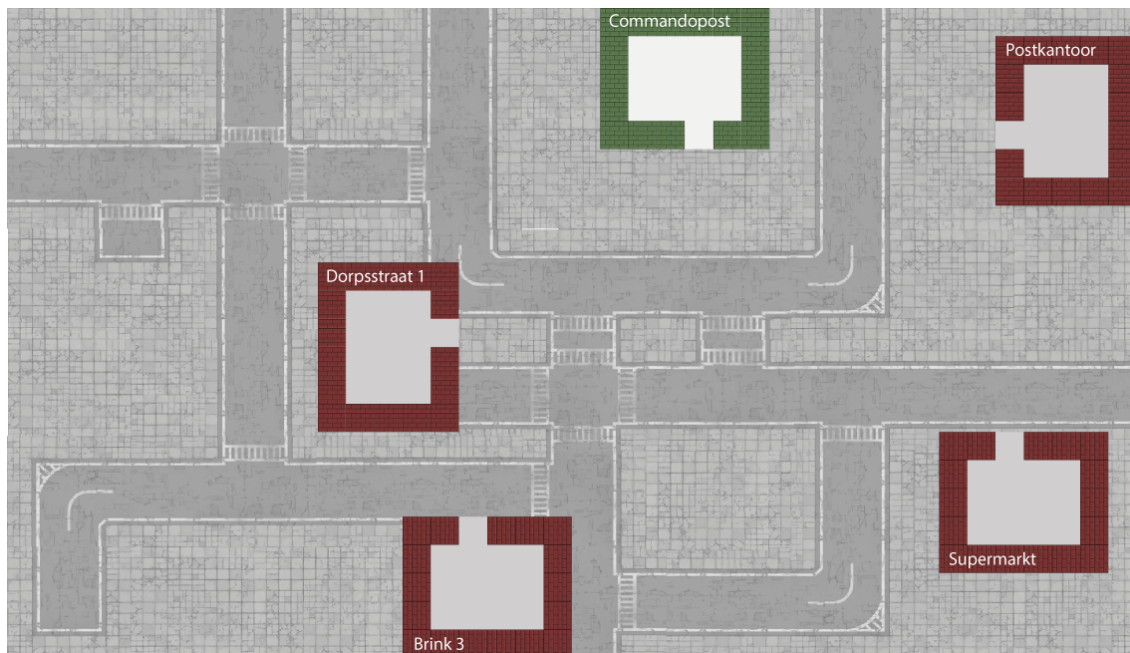


Figure 17: Which of the two messages did you send first to the robot?

is likely to be when an earthquake is approaching.
You may select multiple locations if you wish.

Question 23

Is there a reason for the robot to stay in that spot, do you think?
Please explain your answer.

Question 24

Was there a way to know that an earthquake was imminent?
If so, how?

Question 25

What did you when an earthquake was coming? Please explain your answer.

Appendix E. Implementation of Timing Thresholds for Agent Behavior

The values of the timings were programmed in number of ticks. These were as follows for the following events:

- when the participant ignored the robot’s slowing down navigation due to mud, THEN after 100 ticks (=approx. 10s), HQ issued a reflection question to the participant (see Appendix A)
- When the participant remained standing next to a severely wounded victim (and did not send any messages), THEN after 100 ticks (=approx. 10s), HQ issued a reflection question to the participant (see Appendix A).
- when the participant did not seek shelter for an upcoming earthquake after 50 ticks (=approx. 5s), THEN the robot issued an explanation; IF participant was still not seeking shelter within 5s, THEN the robot issued an assignment; IF still not seeking shelter within 5s, THEN HQ issued an appeal to reflect.
- when the participant is not moving towards the robot (to offer assistance) to open a jammed door; THEN the robot issued an explanation; IF participants still not offered assistance within 1s, THEN the robot issued an assignment; IF still not seeking shelter within 1s, THEN HQ issued an appeal to reflect. (even though the timings of the jammed doors were selected on the basis of pilot participants (see Appendix F), the final timings proved to be on the short side when we conducted the experiment proper).

Appendix F. Determining Timing Thresholds for Agent Behavior

Timers were determined for each cue individually. The experimenter first guessed a time for the robot to produce the cueing behavior. This was tuned in pilot experiments using eight naive intern students, who were one-by-one asked to play the role of human participants. We asked the first participating students afterwards whether it was clear to them that the robot was waiting for the participant to demonstrate the required behavior. If the student indicated that they did not realize this, we increased the waiting time. If the robot commenced walking away while the pilot participant was still busy sending a response, or when they just started with walking towards the robot, then too the waiting time was increased. If, however, the pilot participant pointed out that the robot stood still for a long time, then the waiting time was decreased. Then the new settings were used for the next pilot-participant. The tuned settings after eight pilot students were used for the experiment proper.

Appendix G. Demographics

Table 6: gender-distribution of participants, split by condition.

		sex			Total
		male	female	rather not say	
condition	emergent learning	11	11	0	22
	intentional learning	10	12	1	23

Table 7: age-distribution of participants, split by condition.

		age category							Total
		18-25	26-30	31-35	36-40	41-45	46-50	51-55	
		year	year	year	year	year	year	year	
condition	emergent learning	8	7	4	1	1	0	1	22
	intentional learning	7	7	3	1	2	3	0	23

Table 8: distribution of game-experience of participants, split by condition.

		game experience					Total
		no	some	fair	high	very high	
		experience	experience	experience	experience	experience	
condition	emergent learning	1	12	3	4	2	22
	intentional learning	2	11	5	3	2	23

References

- Ajoudani, A., Zanchettin, A. M., Ivaldi, S., Albu-Schäffer, A., Kosuge, K., & Khatib, O. (2018). Progress and prospects of the human–robot collaboration. *Autonomous Robots*, 42(5), 957–975.
- Anjomshoe, S., Najjar, A., Calvaresi, D., & Främling, K. (2019). Explainable agents and robots: Results from a systematic literature review. In *Proceedings of the 18th International Conference on Autonomous Agents and MultiAgent Systems*, pp. 1078–1088. International Foundation for Autonomous Agents and Multiagent Systems.
- Apperly, I. A., Riggs, K. J., Simpson, A., Chiavarino, C., & Samson, D. (2006). Is belief reasoning automatic?. *Psychological Science*, 17(10), 841–844.
- Barber, D., Davis, L., Nicholson, D., Finkelstein, N., & Chen, J. Y. (2008). The mixed initiative experimental (MIX) testbed for human robot interactions with varied levels of automation. In *Proceedings of the 26th Army Science Conference*, pp. 1–4.
- Bartlett, C. E., & Cooke, N. J. (2015). Human-robot teaming in urban search and rescue. In *Proceedings of the Human Factors and Ergonomics Society Annual Meeting*, Vol. 59, pp. 250–254. SAGE Publications Sage CA: Los Angeles, CA.
- Beans, C. (2018). Can robots make good teammates?. *Proceedings of the National Academy of Sciences*, 115(44), 11106–11108.
- Bolander, T., Dissing, L., & Herrmann, N. (2021). DEL-based epistemic planning for human-robot collaboration: Theory and implementation. In *Proceedings of the International*

Conference on Principles of Knowledge Representation and Reasoning, Vol. 18, pp. 120–129.

- Bradshaw, J. M., Dignum, V., Jonker, C., & Sierhuis, M. (2012). Human-agent-robot teamwork. *IEEE Intelligent Systems*, 27(2), 8–13.
- Cannon-Bowers, J. A., Salas, E., & Converse, S. (1993). Shared mental models in expert team decision making. In *Individual and Group Decision Making: Current Issues*, pp. 221–246. Lawrence Erlbaum Associates, Inc, Hillsdale, NJ, US.
- Colagrossi, M. (2018). What’s the difference between A.I., machine learning, and robotics?. <https://bigthink.com/technology-innovation/whats-the-difference-between-ai-machine-learning-and-robotics/>.
- Correia, F., Mascarenhas, S., Prada, R., Melo, F. S., & Paiva, A. (2018). Group-based Emotions in Teams of Humans and Robots. In *Proceedings of the 2018 ACM/IEEE International Conference on Human-Robot Interaction, HRI '18*, pp. 261–269, New York, NY, USA. Association for Computing Machinery.
- Correia, F., Melo, F. S., & Paiva, A. (2024). When a Robot Is Your Teammate. *Topics in Cognitive Science*, 16(3), 527–553.
- Critchfield, T. S., & Twyman, J. S. (2014). Prospective instructional design: Establishing conditions for emergent learning. *Journal of Cognitive Education and Psychology*, 13(2), 201–217.
- Cross, E. S., & Ramsey, R. (2021). Mind Meets Machine: Towards a Cognitive Science of Human–Machine Interactions. *Trends in Cognitive Sciences*, 25(3), 200–212.
- Cuevas, H. M., Fiore, S. M., Caldwell, B. S., & Strater, L. (2007). Augmenting team cognition in human-automation teams performing in complex operational environments. *Aviation, space, and environmental medicine*, 78(5), B63–B70.
- Cunneen, M., Mullins, M., & Murphy, F. (2019). Autonomous vehicles and embedded artificial intelligence: The challenges of framing machine driving decisions. *Applied Artificial Intelligence*, 33(8), 706–731.
- da Silva, H. H., Rocha, M., Trajano, G., Morales, A. S., Sarkadi, S., & Panisson, A. R. (2024). Distributed Theory of Mind in Multi-Agent Systems. In *16th International Conference on Agents and Artificial Intelligence (ICAART 2024)*. SciTePress.
- Demir, M., McNeese, N. J., & Cooke, N. J. (2018). Dyadic team interaction and shared cognition to inform human-robot teaming. In *Proceedings of the Human Factors and Ergonomics Society Annual Meeting*, Vol. 62, pp. 124–124. SAGE Publications Sage CA: Los Angeles, CA.
- Demir, M., McNeese, N. J., & Cooke, N. J. (2020). Understanding human-robot teams in light of all-human teams: Aspects of team interaction and shared cognition. *International Journal of Human-Computer Studies*, 140, 102436.
- Demir, M., McNeese, N. J., Johnson, C., Gorman, J. C., Grimm, D., & Cooke, N. J. (2019). Effective Team Interaction for Adaptive Training and Situation Awareness in Human-Autonomy Teaming. In *2019 IEEE Conference on Cognitive and Computational Aspects of Situation Management (CogSIMA)*, pp. 122–126. IEEE.

- Devin, S., & Alami, R. (2016). An implemented theory of mind to improve human-robot shared plans execution. In *2016 11th ACM/IEEE International Conference on Human-Robot Interaction (HRI)*, pp. 319–326. IEEE.
- Dissing, L., & Bolander, T. (2020). Implementing Theory of Mind on a Robot Using Dynamic Epistemic Logic.. In *IJCAI*, pp. 1615–1621.
- Driskell, T., Salas, E., & Driskell, J. E. (2018). Teams in extreme environments: Alterations in team development and teamwork. *Human Resource Management Review*, 28(4), 434–449.
- Dunin-Keplicz, B., & Verbrugge, R. (2011). *Teamwork in Multi-Agent Systems: A Formal Approach*. John Wiley & Sons.
- Fagin, R., Halpern, J. Y., Moses, Y., & Vardi, M. Y. (1999). Common knowledge revisited. *Annals of Pure and Applied Logic*, 96(1-3), 89–105.
- Gogas, P., & Papadimitriou, T. (2021). Machine learning in economics and finance. *Computational Economics*, 57, 1–4.
- Groom, V., & Nass, C. (2007). Can robots be teammates?: Benchmarks in human-robot teams. *Interaction Studies*, 8(3), 483–500.
- Harbers, M., Bradshaw, J. M., Johnson, M., Feltovich, P., Van Den Bosch, K., & Meyer, J.-J. (2011). Explanation and coordination in human-agent teams: A study in the BW4T testbed. In *2011 IEEE/WIC/ACM International Conferences on Web Intelligence and Intelligent Agent Technology*, Vol. 3, pp. 17–20. IEEE.
- Hoffman, G. (2019). Evaluating fluency in human-robot collaboration. *IEEE Transactions on Human-Machine Systems*, 49(3), 209–218.
- Holec, R., Hockensmith, M., Broll, J., Wittich, C., Donadio, B., de Visser, E., & Tossell, C. (2020). Autonomous Flight Testbed (AFT): Designing a Flight Simulation System to Explore Future Human-Machine Teaming Concepts. Tech. rep., EasyChair.
- Holstein, K., Aleven, V., & Rummel, N. (2020). A Conceptual Framework for Human-AI Hybrid Adaptivity in Education. In *International Conference on Artificial Intelligence in Education*, pp. 240–254. Springer.
- Johnson, M., Bradshaw, J., Duran, D., Vignati, M., Feltovich, P., Jonker, C., & van Riemsdijk, M. B. (2015). RT4T: A reconfigurable testbed for joint human-agent-robot teamwork. In *Workshop on Human-Robot Teaming, Proceedings of the Tenth Annual ACM/IEEE International Conference on Human-Robot Interaction*, pp. 254–256.
- Johnson, M., & Bradshaw, J. M. (2021). How Interdependence Explains the World of Teamwork. In Lawless, W. F., Llinas, J., Sofge, D. A., & Mittu, R. (Eds.), *Engineering Artificially Intelligent Systems: A Systems Engineering Approach to Realizing Synergistic Capabilities*, Lecture Notes in Computer Science, pp. 122–146. Springer International Publishing, Cham.
- Johnson, M., Bradshaw, J. M., Feltovich, P. J., Jonker, C. M., Van Riemsdijk, M. B., & Sierhuis, M. (2014). Coactive Design: Designing Support for Interdependence in Joint Activity. *Journal of Human-Robot Interaction*, 3(1), 43.
- Johnson, M., & Vera, A. (2019). No AI Is an Island: The Case for Teaming Intelligence. *AI Magazine*, 40(1), 16–28.

- Jonker, C. M., Van Riemsdijk, M. B., Vermeulen, B., & Den Helder, F. (2010). Shared mental models: A Conceptual Analysis. In *In: Coordination, Organization, Institutions and Norms in Multi-Agent Systems at AAMAS2010*. Citeseer.
- Keysar, B., Lin, S., & Barr, D. J. (2003). Limits on theory of mind use in adults. *Cognition*, 89(1), 25–41.
- Kiger, M. E., & Varpio, L. (2020). Thematic analysis of qualitative data: AMEE Guide No. 131. *Medical Teacher*, 42(8), 846–854.
- Klein, G., Woods, D. D., Bradshaw, J. M., Hoffman, R. R., & Feltovich, P. J. (2004). Ten Challenges for Making Automation a "Team Player" in Joint Human-Agent Activity. *IEEE Intelligent Systems*, 19(06), 91–95.
- Lematta, G. J., Coleman, P. B., Bhatti, S. A., Chiou, E. K., McNeese, N. J., Demir, M., & Cooke, N. J. (2019). Developing Human-Robot Team Interdependence in a Synthetic Task Environment. In *Proceedings of the Human Factors and Ergonomics Society Annual Meeting*, Vol. 63, pp. 1503–1507. SAGE Publications Sage CA: Los Angeles, CA.
- Li, H., Oguntola, I., Hughes, D., Lewis, M., & Sycara, K. (2022). Theory of Mind Modeling in Search and Rescue Teams. In *2022 31st IEEE International Conference on Robot and Human Interactive Communication (RO-MAN)*, pp. 483–489. IEEE.
- Liu, Y., & Nejat, G. (2013). Robotic urban search and rescue: A survey from the control perspective. *Journal of Intelligent & Robotic Systems*, 72(2), 147–165.
- Mathieu, J. E., Heffner, T. S., Goodwin, G. F., Salas, E., & Cannon-Bowers, J. A. (2000). The influence of shared mental models on team process and performance.. *Journal of applied psychology*, 85(2), 273.
- McNeese, N. J., Demir, M., Cooke, N. J., & She, M. (2021). Team Situation Awareness and Conflict: A Study of Human–Machine Teaming. *Journal of Cognitive Engineering and Decision Making*, 15(2-3), 83–96.
- National Academies of Sciences, and Medicine, E. (2021). *Human-AI Teaming: State of the Art and Research Needs*. The National Academies Press, Washington, DC.
- Neerincx, M. A. (2007). Modelling cognitive and affective load for the design of human-machine collaboration. In *Engineering Psychology and Cognitive Ergonomics: 7th International Conference, EPCE 2007, Held as Part of HCI International 2007, Beijing, China, July 22-27, 2007. Proceedings 7*, pp. 568–574. Springer.
- Nikolaidis, S., Hsu, D., & Srinivasa, S. (2017). Human-robot mutual adaptation in collaborative tasks: Models and experiments. *The International Journal of Robotics Research*, 36(5-7), 618–634.
- O'Neill, T., McNeese, N., Barron, A., & Schelble, B. (2022). Human–Autonomy Teaming: A Review and Analysis of the Empirical Literature. *Human Factors*, 64(5), 904–938.
- Patterson, R. E., Pierce, B. J., Bell, H. H., & Klein, G. (2010). Implicit learning, tacit knowledge, expertise development, and naturalistic decision making. *Journal of Cognitive Engineering and Decision Making*, 4(4), 289–303.

- Peeters, M. M. M., van Diggelen, J., van den Bosch, K., Bronkhorst, A., Neerincx, M. A., Schraagen, J. M., & Raaijmakers, S. (2021). Hybrid collective intelligence in a human–AI society. *AI & SOCIETY*, 36(1), 217–238.
- Premack, D., & Woodruff, G. (1978). Does the chimpanzee have a theory of mind?. *Behavioral and Brain Sciences*, 1(4), 515–526.
- Salas, E., Shuffler, M. L., Thayer, A. L., Bedwell, W. L., & Lazzara, E. H. (2015). Understanding and improving teamwork in organizations: A scientifically based practical guide. *Human resource management*, 54(4), 599–622.
- Schmidt, R. A., & Bjork, R. A. (1992). New conceptualizations of practice: Common principles in three paradigms suggest new concepts for training. *Psychological science*, 3(4), 207–218.
- Schoonderwoerd, T. A., Van Zoelen, E. M., Van Den Bosch, K., & Neerincx, M. A. (2022). Design patterns for human-AI co-learning: A wizard-of-Oz evaluation in an urban-search-and-rescue task. *International Journal of Human-Computer Studies*, 164, 102831.
- Sen, S., Kremer, L., & Buxbaum, H. (2019). Human-Robot Interaction in Health Care Automation. In *Digital Human Modeling and Applications in Health, Safety, Ergonomics and Risk Management. Healthcare Applications: 10th International Conference, DHM 2019, Held as Part of the 21st HCI International Conference, HCII 2019, Orlando, FL, USA, July 26–31, 2019, Proceedings, Part II 21*, pp. 156–168. Springer.
- Stalnaker, R. (2002). Common ground. *Linguistics and philosophy*, 25(5/6), 701–721.
- Sycara, K., & Sukthankar, G. (2006). Literature review of teamwork models. *Robotics Institute, Carnegie Mellon University*, 31(31), 1–31.
- Van Den Bosch, K., Schoonderwoerd, T., Blankendaal, R., & Neerincx, M. (2019). Six Challenges for Human-AI Co-learning. In Sottilare, R. A., & Schwarz, J. (Eds.), *Adaptive Instructional Systems*, Vol. 11597, pp. 572–589. Springer International Publishing, Cham.
- Van Diggelen, J., Neerincx, M., Peeters, M., & Schraagen, J. M. (2018). Developing effective and resilient human-agent teamwork using team design patterns. *IEEE intelligent systems*, 34(2), 15–24.
- Van Zoelen, E. M., Van Den Bosch, K., & Neerincx, M. (2021a). Becoming Team Members: Identifying Interaction Patterns of Mutual Adaptation for Human-Robot Co-Learning. *Frontiers in Robotics and AI*, 8.
- Van Zoelen, E. M., Van Den Bosch, K., & Neerincx, M. (2021b). Human-Robot Co-Learning for Fluent Collaborations. In *Companion of the 2021 ACM/IEEE International Conference on Human-Robot Interaction*, pp. 574–576.
- Verbrugge, R. (2009). Logic and social cognition: The facts matter, and so do computational models. *Journal of Philosophical Logic*, 38(6), 649–680.
- Verbrugge, R., & Mol, L. (2008). Learning to apply theory of mind. *Journal of Logic, Language and Information*, 17, 489–511.
- Winfield, A. F. (2018). Experiments in artificial theory of mind: From safety to story-telling. *Frontiers in Robotics and AI*, 5, 75.