



Design patterns for human-AI co-learning: A wizard-of-Oz evaluation in an urban-search-and-rescue task

Tjeerd A.J. Schoonderwoerd^a, Emma M. van Zoelen^{*,a,b}, Karel van den Bosch^a, Mark A. Neerinx^{a,b}

^a TNO, Soesterberg, Kampweg 55, the Netherlands

^b Delft University of Technology, Delft, Mekelweg 5, the Netherlands

ARTICLE INFO

Keywords:

Human-AI co-learning
Human-AI collaboration
Design patterns
Learning design patterns
Urban-search-and-rescue
Wizard-of-Oz study

ABSTRACT

The rapid advancement of technology empowered by artificial intelligence is believed to intensify the collaboration between humans and AI as team partners. Successful collaboration requires partners to learn about each other and about the task. This human-AI co-learning can be achieved by presenting situations that enable partners to share knowledge and experiences. In this paper we describe the development and implementation of a task context and procedures for studying co-learning. More specifically, we designed specific sequences of interactions that aim to initiate and facilitate the co-learning process. The effects of these interventions on learning were evaluated in an experiment, using a simplified virtual urban-search-and-rescue task for a human-robot team. The human participants performed a victim rescue- and evacuation mission in collaboration with a *wizard-of-Oz* (i.e., a confederate of the experimenter who executed the robot-behavior consistent with an ontology-based AI-model). The designed interaction sequences, formulated as Learning Design Patterns (LDPs), were intended to bring about co-learning. Results show that LDPs support the humans understanding and awareness of their robot partner and of the teamwork. No effects were found on collaboration fluency, nor on team performance. Results are used to discuss the importance of co-learning, the challenges of designing human-AI team tasks for research into this phenomenon, and the conditions under which co-learning is likely to be successful. The study contributes to our understanding of how humans learn with and from AI-partners, and our propositions for designing intentional learning (LDPs) provide directions for applications in future human-AI teams.

1. Introduction

The increasing advancements in the development and deployment of technology utilizing artificial intelligence are changing the way individuals and teams learn and perform their tasks. It is believed that in the future, humans and intelligent machines will operate more jointly, as hybrid teams (e.g., Li et al., 2015; Peeters et al., 2020; Woods et al., 2004). To enable a team to harmonize its work processes, it is important to be familiar with team members social, cognitive, affective and physical qualities (Demir et al., 2020; Ososky et al., 2012). The development of a hybrid team therefore requires the team to be frequently involved in situations that enable and support partners to learn about the task and about each other. In addition, it requires collaborative learning: situations in which both humans and agents learn how the performance of the team depends upon their own role, upon the role of

the other members in the team, and upon the interdependencies between them Stout et al. (2017). These situations should enable partners to learn about a wide array of characteristics of others, such as a team members objectives, skills, its (work) history of relevant past experiences; its inclination to request or offer assistance; its motivation to contribute to the teams objectives; and many more properties. For such situations, we use the term *co-learning* rather than just learning, because it involves learning from interactions, and has the explicit objective of learning together in order to improve team functioning and performance. Co-learning supports a team to develop from a collection of separate team members into a coordinated expert team (Salas et al., 1997).

In contrast to human-human teams, the members of a hybrid team have different information processing systems, they bring in different knowledge about the task and domain, and do not naturally and

* Corresponding author.

E-mail address: e.m.vanzoelen@tudelft.nl (E.M. Zoelen).

<https://doi.org/10.1016/j.ijhcs.2022.102831>

Received 7 July 2021; Received in revised form 22 December 2021; Accepted 13 February 2022

Available online 1 April 2022

1071-5819/© 2022 The Author(s). Published by Elsevier Ltd. This is an open access article under the CC BY-NC-ND license (<http://creativecommons.org/licenses/by-nc-nd/4.0/>).

automatically share a language to communicate about their knowledge, intentions and plans. Yet, despite these differences, humans and agents need to develop and gradually refine the knowledge, understanding, and skills that are needed for successful cooperation as a hybrid team. To support the team in this co-learning process, methods are needed that enable human and AI team members to share their knowledge and experiences with each other, while accommodating to their inherent differences. Such methods should ideally be generic in nature, allowing its use in various situations. Moreover, proposed methods should be evaluated in a real or simulated task environment to determine their effects on team functioning and performance.

In this paper, we discuss how to design human-AI co-learning. Based on principles from the literature on team learning and human-AI collaboration, we designed a set of sequenced interactions intended to initiate learning of specific objectives. We introduce the term Learning Design Patterns (LDPs) for this. An LDP should be fit for recurrent use in a variety of situations that require team partners to learn about the task and about the team. In an experiment for a human-AI team, we developed two LDPs and empirically evaluated the effects on learning in the human. The team consisted of one human participant and one robot AI, who jointly performed an Urban-Search-And-Rescue (USAR) mission in a simulated environment. The robot was controlled by a *wizard-of-Oz*-experimenter, a person behind the scenes (Riek, 2012). This technique allows studying human-robot interaction without the need of computationally modeling all the required prerequisite competencies of the robot, like sensing the environment and communicating in natural language. That is, existing computational robot models lack the functionality and flexibility for studying how members will be able to learn within a future human-robot team, as these models do not yet sufficiently incorporate the principles of interdependence and autonomy (Lematta et al., 2019). In this study we use the method of a restricted-perception wizard-of-Oz (WOz), that has been advocated for the study of designs for strategies in human-robot interaction research (Sequeira et al., 2016). Half of the human-robot teams engaged in the LDPs; the other teams did not. We investigated the effects of LDPs on task critical knowledge and situational team awareness (Stanton et al., 2017), both being critically important for coordinated team operation, and the effects on the teams overall performance.

2. Theoretical background

There is a rapidly increasing body of research in human-AI teaming and human-robot collaboration (Ajoudani et al., 2018). Application areas have mostly been safety-critical contexts (Bradshaw et al., 2003; Kruijff et al., 2014) and manufacturing (Matheson et al., 2019). Recently it has extended to other domains, for example to healthcare (Buxbaum et al., 2019). Many studies address the utilization of the different strengths and weaknesses of human and artificial intelligence. In order for a hybrid team to make use of the different capabilities of the AI technology and the human, members need to be able to collaborate fluently. Demands for creating successful human-AI collaborations are: (1) conditions in which all partners come to recognize and acknowledge their respective capabilities; (2) a shared understanding of how to exploit complementary strengths to the benefit of the team; and (3) a method for establishing adjusted and new work agreements based on the team partners' progressive insights (Mioch et al., 2018).

2.1. Co-learning

Developing fluent collaborations is a challenge, even in human-only teams. The development of a teams competency is brought about by interrelated processes, ranging from team partners temporarily coordinating their activities in response to local task circumstances in the short term, to fine-tuning their actions to accommodate variations that may re-occur in the task context in the long-term. This is all co-learning (van den Bosch et al., 2019): a process in which collaborating partners adapt

to each other and learn together over time. It is key that such learning does not happen separately, but through collaborative interactions that enable humans and AI to discover and learn about the task, themselves, and their team partners. Moreover, to be able to cope with dynamic environments, learning should take place in situations that closely resemble the actual work environment. Learning should not be limited to formal training, but should continue during the lifetime of a teams operation, embedded in on-the-job work. Learning always takes place, with every new exercise or performance of a team (e.g., Mitchell et al. (2018)).

Co-learning in human-AI teams is related to collaborative learning within humans-only teams (Dillenbourg et al., 1996), but not the same, considering that human and AI team members have different kinds of mental models, embodiments, and ways of learning. A graphical overview of co-learning can be seen in Fig. 1, which shows the interactions between team members and their environment, as well as the growth of their individual mental models and their shared mental model. Co-learning has been identified as important for successful human-AI teamwork, and its components have been conceptually investigated (Holstein et al., 2020; van den Bosch et al., 2019; Wenskovitch and North, 2020). In addition to this conceptual work, there is a need for empirical research into the design of co-learning for human-AI teams, and into its effects on team processes and team performance.

Co-learning may occur implicitly, while partners jointly perform the task. From experience, they learn what sequences, or patterns of interaction (e.g., explaining certain actions, or requesting assistance for a particular task) contribute to the teams mission, and what sequences or patterns are not successful (this relates to co-adaptation, as in Nikolaidis et al. (2017)). Partners may become consciously aware which particular patterns are successful, but explicit awareness is in itself not necessary for partners to learn and apply this knowledge (Patterson et al. (2010)). In fact, such learning from experience often remains tacit (Reber et al., 2019). In contrast, co-learning may also take place intentionally, in situations purposely designed to elicit interactions that enable team members to learn, and to become explicitly aware of what has been learned. Such formalization of what has been learned supports partners to sustain successful interactions beyond the training context.

2.2. Learning design patterns

In our study, we focus on co-learning through explicitly designed interactions in which prescribed learning activities enable team members to improve, correct and extend their mental models. We describe the learning interactions in terms of Design Patterns. Design Patterns are used to describe a solution to a generic or recurring design problem within a particular context (Alexander, 1977; Van Welie et al., 2001). In our case, the design problem consists of the learning that needs to take place between two members of a team. We compose Learning Design Patterns to specify the interactions that take place between these team members to facilitate co-learning. The LDPs can be seen as an extension of Team Design Patterns (van Diggelen and Johnson, 2019), which have been specifically developed to guide human-AI teams in different configurations. Our LDPs aim to optimize the co-learning process, to improve long-term team performance, in current as well as future tasks and contexts.

3. Design of human-AI team context

To utilize the different strengths and weaknesses of human and artificial intelligence, a hybrid team should be designed for interdependence in human-AI relationships (van den Bosch et al., 2019). Within such a team, an AI-robot needs to be able to coordinate its activities with that of other members of the team; it should be able to provide or request help, and it should be able to collaborate with another team member on the same task. A true hybrid human-AI team therefore sets demands with respect to observability, predictability, directability, and explainability

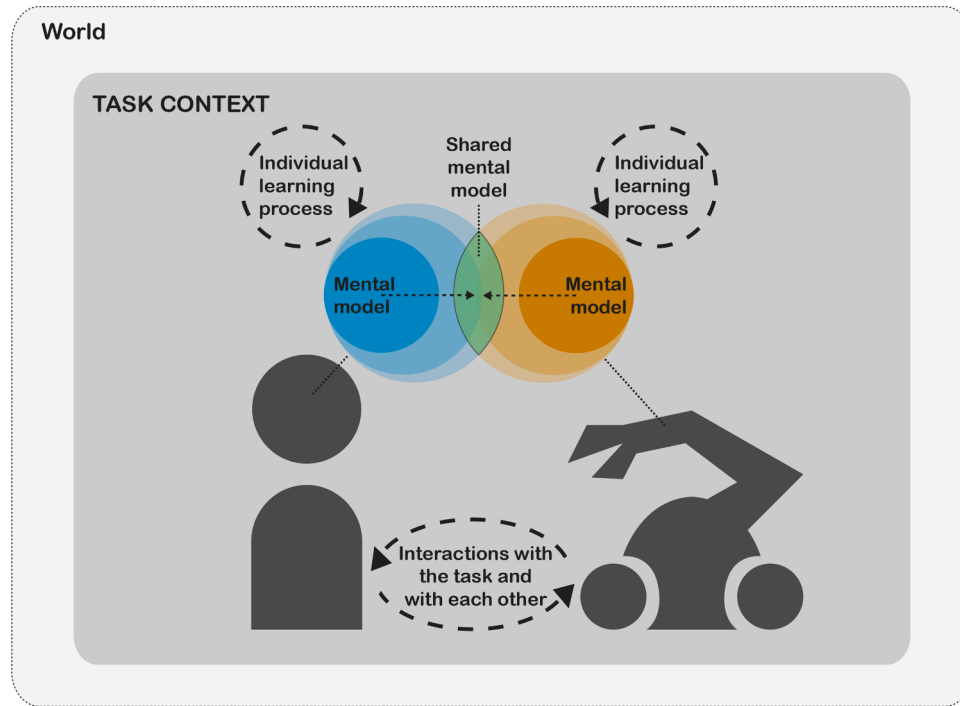


Fig. 1. Co-learning in a human-AI team.

of its team members (Johnson et al., 2014a; Klein et al., 2004; Peeters et al., 2020).

A context that meets the above requirements involves dependencies among the members of a team. Johnson and colleagues (Johnson and Bradshaw, 2021; Johnson et al., 2014a) define dependency in terms of capacity and relationships. *Capacity* refers to the knowledge, skills, abilities, and resources that a team member requires to competently perform an activity individually. Dependency exists when a member lacks a required capacity to competently perform an activity in a given context. *Relationships* refer to the ability to regulate one's own behavior in response to the needs of another team member, and to the requirements of the team's task. This may pertain, for example, to: synchronizing actions, delegating or taking over tasks, and issuing authorizations to permit or prohibit various actions. Dependency exists when a member cannot perform a particular task without the help of a team member (e.g., together carrying a voluminous and heavy object in order to move it), or when a member can perform a task much better and quicker when supported by a team member (e.g., an idle member taking over a task from a very busy member reduces the work load of the team and speeds up completion time).

Thus, designing a context for studying human-AI co-learning requires dependency between tasks, and interdependency between team members. It compels team members to support one another in normal and unexpected situations, and to best utilize the strengths of each.

3.1. A team task for studying human-AI co-learning

It has been advocated that Urban-Search-And-Rescue (USAR) is an appropriate domain for studying how learning of human and AI team members may be investigated and supported (Lematta et al., 2019). In future USAR teams, robots are expected to fulfill cognitive task functions in victim identification that were previously carried out by people, including reasoning with mental models (Sreedharan and Kambhampati, 2018), communicating in natural language (Feng et al., 2018), and providing explanations (Chakraborti et al., 2017). We developed a computer simulation of an USAR task, in which a human and an agent (representing a robot, controlled by a wizard) have to jointly perform

the search and evacuation of victims from an incident area (see Fig. 2).

There are a number of buildings in the area that has been hit by an earthquake. Each may contain one or more victims. The walls of the buildings are shown as colored squares. These buildings may be damaged by the earthquake, and each building may contain one or more victims. Victims may be unhurt, wounded or dead. The team has to localize all victims in the buildings, assess their condition, and bring them to the command post. Dependencies have been built in the task. For example, wounded victims need to be treated first before they can be brought to the command post. Also, damaged buildings cannot be entered unless the debris blocking the entrance has been removed. Interdependencies between team members have been implemented by assigning complementary capacities to the human and the agent. For example, only the robot can assess whether or not a building has been damaged by the earthquake. Furthermore, the robot can remove debris, but the human cannot. In contrast, the human can treat victims, but the robot cannot. However, they can both carry victims from the buildings to the command post. By using complementary capacities, the team requires collaboration to complete the task. The human can send commands to the robot using template sentences (e.g., "Free entrance of building...") via a chat-box that is displayed next to the task environment.

In our experiment, we employ a WOZ-paradigm in which the agent is controlled by a confederate researcher. To determine the agent's behavior in the scenario, the wizard strictly followed a behavior protocol that was created based on the knowledge and behavior model of the agent (see Section 4). This protocol dictated the order of task actions that the agent should perform (e.g., move to closest building, inspect the building, clear the entrance), and the behavior of the agent in response to events or actions of the human team member (e.g., the human team member sending a command, or going into a building that the robot was already navigating towards).

To introduce a controlled need for co-learning, the robot's model (see Section 4) lacked certain knowledge elements that are of critical importance for efficient task execution. The implications of this knowledge deficiency appear approximately halfway the first run of the scenario, when the area is suddenly hit by a second earthquake. In the

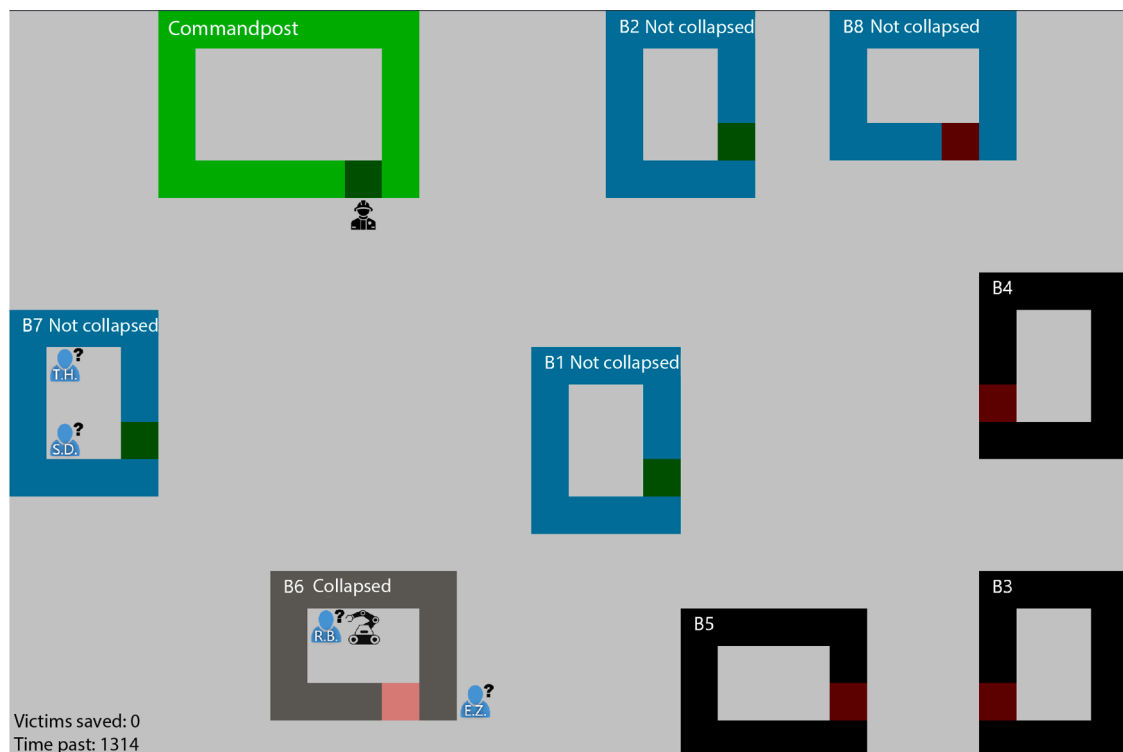


Fig. 2. Screenshot of the human-robot team performing an urban-search-and-rescue task. The avatar with the helmet is of the human participant; the agent's avatar is in building B6. The avatars with initials represent victims, and the question mark means that their condition has not yet been assessed. The color of a building indicates its status, with grey indicating: collapsed; and blue: not collapsed. The green building is the command post. (For interpretation of the references to colour in this figure legend, the reader is referred to the web version of this article.)

instruction, before the start of the experiment, the human participant was informed that if such a second earthquake would occur, the status of already examined buildings would expire. A re-examination of the buildings in the vicinity of the earthquake would then be needed. However, this knowledge was not part of the robot's model, and thus the behavior protocol for the wizard controlling the robot did not include actions based on this knowledge. The human participant was not informed about the robot being unfamiliar with this procedure. As a consequence, soon after the second earthquake, the human participant was confronted with unexpected behavior of the robot. Participants had to sort out for themselves how to proceed and to complete the run. Then, for half of the teams, two Learning Design Patterns were initiated (see Section 5 for details). The other half of the teams did not engage in these LDPs. Then, a second run was administered to all teams, which again included an earthquake midway the scenario. Measures of collaboration, mutual understanding, and performance (see Section 6.2.2) were used to test the effects of the Learning Design Patterns.

4. Knowledge and behavior model of the robot

Although the WOZ-paradigm allows a human controller to freely determine the agent behavior, we chose to explicitly model the knowledge and behavior of the robot. We deem this necessary for maintaining ecological validity of research into human-AI teams, because humans should have the impression that they are working with an AI-driven teammate and not with another human being. In order to simulate the AI in the robot team member, we made a formal representation of the robot's knowledge, and a behavioral model that links knowledge to specific robot behaviors. The behavioral model was used to construct the protocol for the wizard, which dictated how the agent should be controlled during the USAR task.

The knowledge of the robot was represented by an ontology-based model containing concepts related to the task (e.g., goals,

requirements) and to team members (e.g., capabilities, intentions). This knowledge was determined prior to the experiment. The model consists of a domain ontology that represents a relational network of concepts (classes or properties) related to the USAR task (e.g., Goal, Agent, Role, BuildingStatus, Location, Robot, VictimStatus, Earthquake). This domain ontology was built on top of the SUMO upper ontology (Niles and Pease, 2001), and extended with management concepts (MSPM; Cheah (2008)) and task world models (Van Welie et al., 1998). The final ontology consisted of 53 classes (19 of which are object properties) and 248 axioms. As an example, Human and Robot are modelled as sub-classes of CognitiveAgent (i.e., an agent with responsibilities and with the ability to reason; (Niles and Pease, 2001)), that have an Identifier, Skill(s), and a Location. The concepts Explorer and RescueWorker are Roles that can be enacted by CognitiveAgents, are associated to a Task, and require at least one Skill.

Concepts from the knowledge model were used to construct a behavior model that defined what actions should be performed by the robot, based on inputs from the environment (i.e., environment state, and commands from the human team member). A goal-driven approach was used, in which the model decomposes a task into sub-goals, and further decomposes it into actions required to achieve a particular goal. The benefit of a goal driven approach is that it -in contrast to machine-learning techniques- allows behavior explanations that are understandable by humans. Arguably, actions from an intelligent system are best understood by humans if they are explained using concepts such as beliefs, intentions, and goals (De Graaf and Malle, 2017; Miller, 2019). Therefore, the robot's behavior model was created using a Goal Hierarchy Tree (GHT, Broekens et al. (2010); Harbers et al. (2010a)), based on the guidelines described in Harbers et al. (2010b). A GHT is a high-level description of the agents reasoning and is based on hierarchical task analysis. We decomposed the USAR task into two high-level goals (find victims in buildings and rescue victims). These were further analyzed into beliefs about world states (e.g., building is collapsed,

victim is mildly injured) and intentions of the robot (e.g., clear building entrance, bring victim to command post).

The robot should respond in a manner that corresponds exactly to its knowledge of the world at that point in time (which is incomplete and even partly incorrect for reasons of the study), its objectives, and to its assigned capacities. Therefore, the protocol for the wizard explicitly linked beliefs about the world state to behavioral intentions through conditional statements (e.g., 'if a building is collapsed, then clear its entrance', and: 'if battery is empty, then send a chat message with current battery level'). The wizard solely made use of this protocol to determine the behavior of the robot.

5. Learning design patterns

We developed two learning design patterns intended to support co-learning. The goal of the first LDP (SitRep LDP) is to support identifying knowledge gaps that team members may have. The objective of the second LDP (Knowledge-rule LDP) is to initiate (inter)actions that enable team partners to learn from other team members.

In our implementation of the USAR task, the need for learning manifests itself directly following the earthquake, as the robot shows behavior that is not expected by the human. The SitRep LDP aims to fulfill this need for learning by providing information that clarifies the robot's behavior; the Knowledge-rule LDP supports the human with teaching the robot critical knowledge about the consequences of the earthquake. The LDPs were developed using a Research through Design approach (Zimmerman et al., 2007): iteratively proposing and evaluating designs. Each design iteration was evaluated by a group of experts. Moreover, user tests were conducted with three students to evaluate the usability and effectiveness of the designs.

The SitRep LDP is intended to be used in situations in which a team member is confounded by the behavior of other team members. Being surprised by behavior of others occurs often, perhaps most typically in beginning teams. The SitRep LDP prescribes the activities that support the confused partner to develop a better understanding of the partner's actions (see Table 1). The SitRep LDP is designed to be applicable in team situations in which confusion and misunderstanding among team partners exists. In the USAR-experiment, the SitRep LDP prescribes the human to request an action report from the agent that is causing the confusion (in our case, the human became confused by the robot moving to an unexpected location after the earthquake). The agent responds by presenting relevant information from its behavior model. In our case, this means showing the beliefs, goals, and intentions from the goal hierarchy tree, that were used to determine the action that caused the confusion. For our experiment, we created a simple graphical interface that enabled participants to obtain this information (see Fig. 3).

The Knowledge-rule LDP is intended to be used in situations that demand learning of task-critical knowledge by a team partner in order to perform adequately. This Knowledge-rule LDP prescribes the activities that supports team members to teach the necessary knowledge to the demanding partner. Again, this LDP is designed to be applicable in team situations where a demand for acquiring task-critical knowledge exists. In our USAR-experiment the LDP prescribes the human to engage in activities that supports the agent to acquire a new knowledge rule that contains the task-critical knowledge (in the form if then) (see Table 2). Fig. 3 shows the graphical interface that we used to enable participants to construct a new rule using concepts from the knowledge model of the robot. The first part of the rule corresponds to the situation in which the robot did not respond correctly (i.e., 'if an earthquake hits during the scenario...'), and was already filled in. Participants completed the 'then...' -part of the knowledge rule by altering concepts (object, property, value) from the robot's knowledge base. After creating a rule, the robot presents feedback information by showing how this knowledge rule will influence its future behavior (e.g., after an earthquake hits, it will first move to buildings for which the status is unknown), by means of a goal-hierarchy tree diagram (inspired by Harbers et al. (2009)).

Table 1
SitRep Learning Design Pattern.

SitRep LDP: Learning from an AI team member by reviewing an action report	
Behavior pattern	An AI team member presents an action report to the human, containing the information that it used to decide to execute the specific action at hand. The human team reads the action report and can click on the different information components of the report to learn more about the background of the AI agent's behavior.
Interaction requirements	The human team member should be able to indicate a particular action carried out by the AI team member. The AI team member should be able to provide a reason why it chose to perform this action.
Positive effect	The human better understands why their AI team member chose to execute a certain action, leading to better future team performance.
Negative effect	The human might draw the wrong conclusions from the presented information, causing misunderstandings. Also, since it takes time and requires team members to interrupt the task, it can affect current team performance negatively.
Use when	When the human wishes to learn more about why an AI team member executed a particular action.
Example	A human and a robot are collaborating to save victims from a disaster area that was hit by an earthquake. The robot checks whether buildings have collapsed and starts to clear blocked doorways. The human diagnoses and treats found victims. At a certain moment, there is an aftershock which hits several buildings. The robot falls silent, and after its battery has been replaced, it continues to check the buildings that it was already planning to check before the aftershock occurred. This behavior violates the procedure and therefore confuses the human. The human asks the AI agent to deliver an action report. The action report contains the current action, the goal that caused this action, the time at which the action execution started, information about the current state of the robot and the environment, and the agent's beliefs that lead to this action (including information that the beliefs are based on). The human reads the information and subsequently understands that the robot does not understand that the aftershock caused additional collapse danger to the buildings. This understanding enables the human to take corrective measures that support appropriate collaboration in future collaborations.
Design rationale	When team members collaborate, they need to understand each others decision making process. It is often argued that common ground is very important to establish trust, and to make AI agents effective team partners (Klein et al., 2004). In the military, a SitRep (Situation Report) is often used to create this common ground between team members (Sorensen and Stanton, 2016). A SitRep is a concise overview of the current situation at hand. It usually contains information about the environment, time and people involved, as well as actions that have been done and will be done in the future. The LDP that we propose is similar to such a SitRep. The information is presented on the basis of progressive disclosure: in terms of high-level beliefs, goals and intentions. It is well known that such information is easily understandable by humans (Dennehy, 1989).

In Tables 1 & 2 we formalize the LDP using the format of Team Design Patterns (Van Diggelen et al., 2018). The presented LDPs should be considered as proto-patterns, meaning that although they are founded on principles and evidence obtained from the literature, their effectiveness at bringing about learning still needs empirical testing. This study contributes to this testing.

By using a WOZ-paradigm to study human-AI co-learning, assumptions must be made concerning the (learning) capabilities of the AI. Considering the LDPs, we assume the AI to be able to explain its actions on request by using symbolic, human-understandable concepts such as objects and properties (LDP 1). Moreover, we assume that the AI makes use of a rule engine to determine its behavior, and that it can adapt its reasoning based on human feedback on its rules (LDP 2). We attempted to make the behavior of the AI partner as realistic as possible by modeling a knowledge base as well as a goal hierarchy tree that could

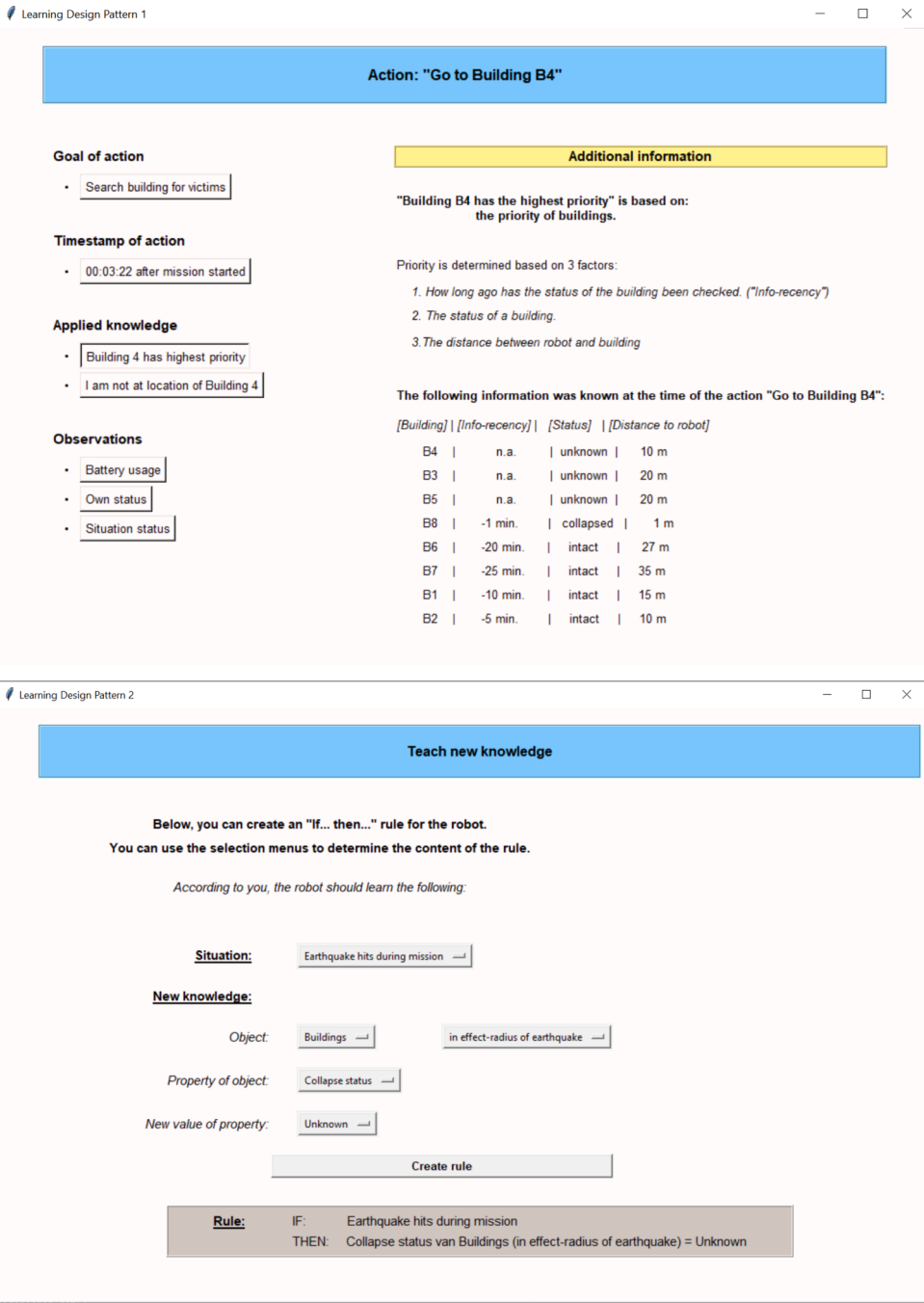


Fig. 3. GUI for the Sitrep LDP (top) and the Knowledge-rule LDP (bottom).

have been part of an actual AI implementation. Instead, we used this knowledge and behavior model to construct the protocol for the wizard, and to create the behavior explanation for LDP 1 and knowledge concepts for LDP 2. The use of this "synthetic" AI was a conscious decision, as it enables us to design task environments and learning interactions for studying co-learning and to explore the complexities of co-learning without having to solve all technical challenges. Of course, the technical challenges posed by human-AI co-learning need also to be solved and require research questions by themselves, but the main focus of the present study is on the human experience of co-learning with an AI partner.

6. Methods

Participants performed two runs of a simulated USAR task in collaboration with a robot. Half of the group was assigned to the experimental condition, in which they performed the Sitrep and Knowledge rule LDPs in between the two tasks. Participants in the control condition did not engage in these learning activities. The robot team partner was controlled by a confederate, following a wizard-of-Oz paradigm (Riek, 2012).

6.1. Participants

Participants were recruited from a participant database. Selection criteria were: 18–45 years of age, and at least a college degree education.

Table 2
Knowledge-rule Learning Design Pattern.

Knowledge-rule LDP: Teach new knowledge to an AI team member	
Behavior pattern	The human implements new knowledge in the AI agent's model in the form of if. then. After implementation, the AI agent presents the effect that this new rule will have on its behavior and action execution.
Interaction requirements	The human should be able to compose a rule that consists of concepts known to the AI agent. The AI team member should be able to identify and present the behavioral changes that would result from implementation of this rule.
Positive effect	The AI gains knowledge and is therefore better able to execute its tasks. Also, as being the implementer of the AI's knowledge, the human team member understands precisely how actions of the AI are brought about.
Negative effect	The human might input a wrong rule when they don't understand the AI agent well enough. This might cause the AI to show unexpected behavior, possibly opposite to the human's intention. Also, since it takes time and requires team members to interrupt their task, it can effect current team performance negatively.
Use when	When a human discovers that the AI agent lacks certain knowledge to properly conduct their task.
Example	A human and a robot are collaborating to save victims from a disaster area that was hit by an earthquake. The robot checks whether buildings have collapsed and starts to clear blocked doorways. The human diagnoses and treats found victims. At a certain moment, there is an aftershock which hits several buildings. The robot falls silent, and after its battery has been replaced, it continues to check the buildings that it was already planning to check before the aftershock occurred. The human understands (through the Sitrep LDP or in another way) that this is because the robot does not understand that the aftershock might further collapse buildings. The human implements into the robot's model the critical knowledge that an aftershock changes the status of nearby buildings to unknown. After implementation, the robot feeds back to the human, step by step, how this will affect its behavior in the future.
Design rationale	When collaborating, there might be situations in which the robot's knowledge is incomplete. The human may infer or recognize this, for example because they received an action report (see Table 1). Humans have the world and task knowledge to infer the knowledge that the AI agent needs to act appropriately in particular situations, and they can define this in a manner that it is re-usable in future other but similar contexts.

A total of 35 participants took part (10 male, 25 female). The median age category was 18–25 years, and the median category of self-rated experience with playing computer games was 1 (no experience). The data from four participants were excluded from analysis; three because of technical complications and one because the participant was unable to complete the task successfully. In total, data from 31 participants were used for analyses. Due to technical complications, task completion time was not registered for five participants.

6.2. Materials

6.2.1. Hardware and software

A simulation of the USAR-task was developed using the Python programming language and a software library called MATRX (multi-agent teaming rapid experimentation)¹. MATRX allows rapid design, simulation, and testing of 2D top-down, grid-based environments in which multiple humans and agents can perform tasks collaboratively. Moreover, it contains a chat window as part of its interface, allowing team members to communicate. We implemented one human rescue worker (controlled by the participant), and one robot explorer agent (controlled by the wizard).

Three laptops were connected through a local network: one server laptop to start and stop the USAR task, and two laptops on which the task was carried out (one for control of the explorer agent, and one for control of the rescue worker agent). The task laptops were equipped with an external keyboard and mouse, which were used to perform actions in the simulation environment (e.g., moving around, treat victims, send commands via the chat). The LDP interfaces ran on the participant laptop, and were coded in Python using the Tkinter² library.

6.2.2. Protocol for robot behavior

Following the guidelines for designing Woz-studies as described in Green et al. (2004), a protocol was created containing instructions for the wizard how to control the robot during the USAR task. The protocol was derived from the goal hierarchy tree (i.e., behavior model) that was constructed for the robot behavior (see Section 4), which links goals (e.g., rescue victims) to beliefs about the world state (e.g., building B2 contains a victim) and intentions (e.g., carry the victim from B2 to the command post) for the robot. In the protocol, we specified what action the wizard should perform based on the current world state. The robot's default behavior policy was: to autonomously move to a building; to determine and report the building's status; and to enter the building and report any victims. When a participant issued a command to the robot, the wizard immediately interrupted the current action and executed the requested command. Upon completion, the wizard resumed actions based on the protocol. The wizard used one and the same protocol for controlling the robot, both in the control and the experimental condition.

In the scenario, the earthquake event was triggered after the robot had inspected four buildings, one of which had to be B1. As the earthquake was programmed to only affect buildings B1 and B2, this ensured that the status of (at least) building B1 was no longer up to date after the earthquake. In the first scenario run, the robot would not respond to the earthquake event at all, and would simply continue its current actions. However, in the second run (for both conditions), the robot responded adequately to the event and directly navigated to the affected buildings (i.e., B1 and B2) to (re)examine their status.

6.2.3. Instruction for participants

Participants received a booklet containing a description of the USAR task. They were given eight short tutorial exercises to get acquainted with the task environment and the controls. An additional information sheet was provided, summarizing the main goal of the task, the capabilities of both agents, and the control scheme. Participants were allowed to consult this sheet at any time. Participants were told that the experiment involved collaboration with an AI-controlled robot. Although in this experiment the robot was actually controlled from another room by a human wizard, the participants were kept unaware of this. They were given the impression that the robot acted autonomously. None of the participants challenged this at any time during the experiment.

6.3. Measures

Two types of measures were performed in the experiment: learning measures to obtain insight into the learning effects of the LDPs on participants, and performance measures to obtain insight into the effects of the LDPs on task performance of the human-AI team. The measurements are summarized in Table 3. Learning processes were measured by using questionnaires and by rating propositions on Likert scales, which were administered after each run (both groups) and after each LDP (experiment group only). The following learning measures were collected:

¹ <https://matrx-software.com/>

² <https://wiki.python.org/moin/TkInter>

Table 3

Measurements taken during the experiment. Quant = Quantitative Data, Qual = Qualitative Data.

Measurement	Data type	Timing
Learning measures		
Accuracy of participant's explanation of overall robot behavior	Qual.	After each run
Accuracy of participant's explanation of the robot behavior after the earthquake	Qual.	After each run
Accuracy, certainty (1–5 Likert scale), and rationale for remembered building choice of the robot after the earthquake	Quant./Qual.	After each run
Expected building choice of the robot after earthquake, and rationale for this choice	Qual.	After each run
Explanation of robots faulty behavior after the earthquake	Qual.	After Sitrep LDP
Certainty of participants knowledge about the robots behavior (1–5 Likert scale)	Qual.	After Sitrep LDP
Correctness of knowledge rule & participants rationale for the rule	Qual.	After Knowledge-Rule LDP
Clarity of the behavior explanation in the Knowledge-rule LDP (1–5 Likert scale) & rationale for the rating	Quant./Qual.	After Knowledge-Rule LDP
Self-rated understanding of the behavior explanation (1–5 Likert scale)	Quant.	After Knowledge-Rule LDP
Performance measures		
Collaboration fluency (avg. rating on 1–5 Likert scale) (Hoffman (2019))	Quant.	After each run
Task duration (# simulation ticks)	Quant.	After each run
Saved victims score (= sum of score per victim, 0=dead, 0.25=severely injured, 0.5=injured, 0.75=slightly injured, 1=uninjured)	Quant.	After each run
Idle time of human agent (#ticks in proportion to task duration)	Quant.	After each run
Number of commands sent by the human to the robot	Quant.	After each run

- what participants learned about the robots behavior (11 items: eight open, three closed questions)
- fluency of collaboration (Dutch translation of Hoffman (2019)). Participants rate 24 statements about their collaboration with the AI team member on a 5-point scale. The test distinguishes five dimensions of collaboration fluency: *human-robot fluency*; *relative contribution of team members*; *trust in the robot*; *positive teammate traits*; *performance improvement over time*; and *working alliance*. Authors of Hoffman (2019) report high reliability of all scales (Cronbachs $\alpha \geq 0.77$).

Experimental group only:

- participant's perception of robot's behavior: after each run, the questionnaire presented open questions about the robots behavior. These questions intended to capture the participant's knowledge and understanding of the robot's behavior. As the objective of an LDP is to improve team member's mental model (see Section 2.2), the questions aim to uncover whether the participant has learned. The complete questionnaire is available online³

Performance measures were obtained by automatically logging data during task execution. The following measures were collected:

- Task duration (shorter completion time means better performance)
- Idle time of the human agent (shorter idle time means better performance)
- The number of commands sent by the human agent (fewer commands means better performance)

- The number and health status of saved victims (higher score means better performance)

6.4. Procedure

The experiment took place in a large conference room, in full adherence to the COVID-19-measures issued by the Dutch government. Participants were randomly assigned to the control group or experimental group. Participants sat behind a table with a laptop, external keyboard and mouse, and the instruction and questionnaire booklets. The wizard controlling the robot was hidden from participants in a side room.

The experiment leader welcomed participants and provided a brief introduction, then took place at the diagonal other side of the table, behind the server laptop. Fig. 4 shows the flow of the experiment. Note that participants from the control group performed both task runs in succession, while participants from the experimental group performed both types of LDPs in between the two runs.

7. Results

Repeated-measures analyses of variance were used to test the within-subjects effects of Run (first vs second) and the between-subjects effects of Group (LDPs versus no-LDPs). Table 4 shows the results of the analyses.

Both groups performed better in the second run. They were quicker and achieved a higher victim score (i.e., more victims are saved, and saved victims were in a healthier condition). The finding that human agents (controlled by participants) were less idle during the second run shows that the human participant was working more efficiently.

Inspection of the interaction between Group and Run for Saved Victims reveals that groups differed on the first run, $t(14) = -2.38$, $p = 0.024$, but not on the second run, $t(14) = 0.47$, $p = 0.644$. A difference between groups on the first run is unexpected, as in that run, both groups were provided with exactly the same task information and instructions. Since no outliers in the saved-victim scores were found in both groups, we regard this finding as incidental. We found no between-groups effects on the number of commands issued by participants.

Groups rated the collaboration fluency within the team equally, and this was not affected by Run. Interestingly, participants rated collaboration fluency relatively high in general (i.e., many rated 4 out of 5) and there was little variation in ratings within and between groups. In the second run, participants rated fluency of collaboration higher than in the first run. Additional exploratory analyses on the subscales of the test showed similar results as for the full scale.

Correlation analyses between objective and subjective performance measures were conducted. Interestingly, for the first run, no statistically significant correlation between fluency ratings and any of the objective performance measures was found. This suggests that performing better as a team does not necessarily imply that people also perceive the collaboration as to be better. For the second run, significant correlations between subjective and objective performance were found for the LDP-group only. A positive correlation between total task duration and fluency ratings was found ($r = 0.71$, $p < 0.001$), showing that, for the LDP-group in the second run, longer task durations were associated with higher ratings on collaboration fluency. Moreover, a higher task duration was accompanied by a higher number of commands ($r = 0.66$, $p < 0.001$), and the number of commands was not correlated with collaboration fluency. This suggests that number-of-commands-issued does not account for the correlation between task duration and fluency scores. As we will argue in Section 7.1, the LDP group paid more attention to the behavior of the robot during the second run. Together this suggests that more time for interaction with the robot means more opportunities to pay attention to the teamwork, resulting in a higher appreciation of the collaboration.

³ Schoonderwoerd, Tjeerd (2021), Questionnaire used in human-AI co-learning experiment, Mendeley Data, V1, doi: 10.17632/7ksfjsgsb2.1.

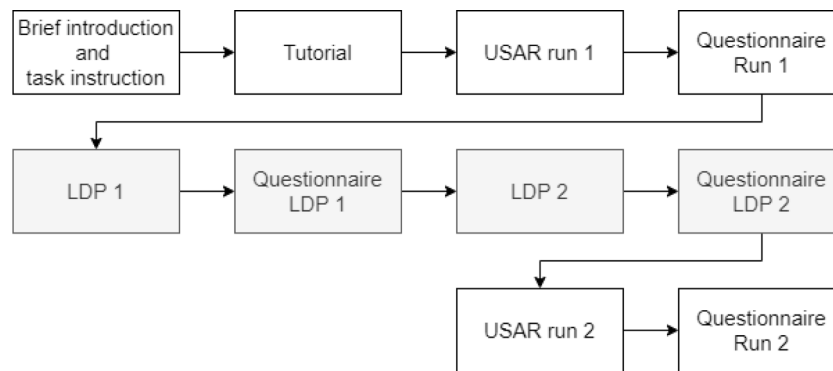


Fig. 4. Flow diagram of the experiment procedure. Grey boxes indicate activities that are exclusively performed by participants from the experimental group.

Table 4

Results of five RM-GLMs to determine the effects within and between groups and runs on each performance metric.

Variable	Means in Run 1 (Control; LDP)	Means in Run 2 (Control; LDP)	Main effect of Group (between-subjects)	Main effect of Run (within-subjects)	Interaction effect (Group x Run)
Task duration (ticks)	4014; 3236	3068; 2731	$F(1,24) = 4.64, p = 0.042^*$	$F(1,24) = 31.41, p < 0.001^*$	$F(1,24) = 2.89, p = 0.102$
Saved victims score	6.70; 7.27	8.05; 7.95	$F(1,29) = 1.60, p = 0.216$	$F(1,29) = 67.00, p < 0.001^*$	$F(1,29) = 7.08, p = 0.013^*$
Idle time (% of duration)	8.00; 7.44	7.33; 7.25	$F(1,24) = 3.94, p = 0.059$	$F(1,24) = 16.56, p < 0.001^*$	$F(1,24) = 3.21, p = 0.086$
Number of commands	9.80; 7.81	7.20; 6.94	$F(1,29) = 0.62, p = 0.439$	$F(1,29) = 2.56, p = 0.120$	$F(1,29) = 0.63, p = 0.434$
Collaboration fluency	3.70; 3.79	3.88; 3.95	$F(1,29) = 0.41, p = 0.528$	$F(1,29) = 6.07, p = 0.020^*$	$F(1,29) = 0.01, p = 0.911$

* Significant at $\alpha = 0.05$.

7.1. Learning results

The open questions after each run were intended to capture the participant's knowledge and understanding of the robot's behavior. The first two authors analyzed the free text responses in the questionnaire. We started with independent open coding of responses to all open questions from three randomly assigned participants. Keywords were then compared and discussed, in order to develop a closed coding scheme for further analysis. This scheme was then used by both evaluators individually to code responses to the open questions from all 31 participants. Then both evaluators compared their outcomes and resolved any differences in assigned keywords through discussion. The final step in the qualitative analysis was to sum the occurrences of keywords for each question per group and per run, in order to obtain a general overview of the responses.

7.1.1. Participants' responses to the task runs

Table 5 shows the results of the analysis of answers on the questions asked after completing each run. After the first run, the frequencies of keywords show only marginal differences between groups. This is expected as participants were randomly assigned to a group and the between-groups manipulation takes place after the first run. After the second run, participants from both groups were positive about the

Table 5

Frequency of keywords in the qualitative analysis of answers on open questions after each run from the control group ($n = 15$) and LDP group ($n = 16$)*.

Measurement	Run	Control group	LDP group
Opinion on overall robot behavior	1	Good (10) Autonomous (3)	Good (13) Autonomous (6)
	2	Good (10) Autonomous (3) Better (2)	Good (5) Better (9)
Explanation of overall robot behavior	1	Fixed task list (9) Algorithm (3)	Fixed task list (7) Algorithm (6)
	2	Fixed task list (6) Algorithm (3)	Fixed task list (4) Addition of rule (4)
Explanation of robot behavior after earthquake	1	Continued with task (3)	Continued with task (6)
	2	Inspect collapsed buildings (6)	Inspect nearby collapsed building (16)
Rationale for remembered building choice of the robot after the earthquake	1	Building in earthquake radius (4)	Building in earthquake radius (4)
	2	Building in earthquake radius (5)	Building status changed (11) Building in earthquake radius (7) Robot is close to building (9)
Rationale for the expected building choice of the robot after the earthquake	1	Building in earthquake radius (10) Building is close to robot (4)	Building in earthquake radius (9) Robot is close to building (5)
	2	Building in earthquake radius (11) Building is close to robot (3)	Building in earthquake radius (8) Robot is close to building (9)

* Results deemed salient by both raters are displayed in bold text.

behavior of the robot during the task in both runs, and some especially praised the autonomy of the robot. Interestingly, more participants from the LDP group stated that the robot had improved its behavior (note that the robot actually behaved the same for the control and LDP group in run 1 and 2). We think that engaging in learning activities may have caused participants of the LDP group to focus on the robot, which may have increased their appreciation of its performance.

Of interest is the question how participants perceive the robot after acquiring their first experiences. This was examined by analyzing responses at the end of run 1. Results showed considerable variation among participants in their views. Most participants believed that the robot acted according to a fixed task list. Some participants described the robot as obedient, goal-oriented, or careful. This suggests that participants did not expect adaptive behavior from the robot.

Participants were asked what might have caused the unexpected behavior of the robot immediately following the earthquake. After the

first run, approximately one-third of all participants (of both groups, as at that point there were no between-groups differences) provided a correct explanation for this behavior (i.e., the robot ignored the earthquake). Other participants stated to be unsure about the cause, or indicate to have no idea. After the second run, there was no change in the control group: still one-third of the participants of the control group (6 out of 15) were able to give the correct explanation. There was a significant change for participants of the LDP group however: all participants provided the correct explanation. This suggests that participants in the LDP group had learned about the robots behavior, while participants in the control group did not.

To further explore the suggestion that participants of the LDP-group were able to substantially improve their awareness and understanding of the robot, it was investigated how well participants were able to recall to what building the robot went directly following the earthquake. After run 1, only one third of participants in each group recalled the building correctly. After run 2, the recall performance was as follows: in the control group, 12 out of 15 participants recalled the correct building. In the LDP-group, the building was correctly recalled by 15 out of 16 participants. The difference between runs was significant ($F(1,29) = 36.49, p < 0.01$), but the difference between groups was not. In addition, we also asked participants how certain they felt about their recall. As shown in Figure 5, both groups felt equally certain in the first run, but more certain on the second run ($F(1,28) = 49.52, p < 0.01$); and the LDP-group felt more certain than the control group ($t(15) = 3.10, p < 0.05$).

Another indication for the notion that participants of the LDP group developed a better understanding and awareness of the robot's behavior was that they were more specific in their explanation of the robot's behavior and the specific circumstances at that time. While the control group mostly used general terms (e.g., building is hit by earthquake, its part of its task list, or: that building has a high risk), the LDP group used more specific terms, apparently taken from the Sitrep LDP (e.g., building status has recently changed, building is close to robot). The higher level of detail in answers of the LDP-participants was also observed when asked to explain the expected robot behavior after the second run.

7.1.2. Participants' Responses to the learning activities

Table 6 shows the keywords and their frequencies obtained from participants answers to the questions after completing each LDP (see Fig. 4). As Table 6 shows, the majority of participants were able to correctly explain the faulty behavior of the robot in their own words. Furthermore they were fairly certain of their explanation, as indicated by the Likert-scale scores ($M = 4.00, SD = 0.73$). Moreover, nearly all participants understood that the robots knowledge model was inaccurate, although only 25% of all participants could point out the fault (i.e.,

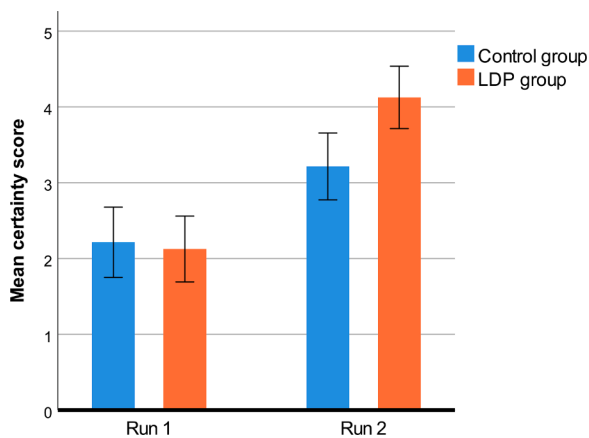


Fig. 5. Mean rating (5-point Likert scale) of participants regarding their certainty about the recalled location of the robot directly following the earthquake. Error bars indicate 95% confidence intervals.

Table 6

Frequency of keywords used to describe answers from participants from the LDP group ($n = 16$) in response to the open questions concerning the Sitrep LDP and Knowledge-rule LDP.

Measurement	LDP	Keywords in LDP group
Explanation of robots faulty behavior after the earthquake	Sitrep	Robot is close to building (11) Building status is unknown (10)
Observed accuracy of knowledge model of the robot, and preferred adjustment to the robot	Sitrep	Inaccurate (15) Behavior agreement (11) Knowledge correction (5)
Correctness of the knowledge rule	Knowledge-rule	Correct (13) Incorrect (3)
Rationale for the knowledge rule	Knowledge-rule	Robot will go to buildings hit by earthquake (12)
Rationale for the clarity rating of the behavior explanation	Knowledge-rule	Step-by-step (8) Visual (4)

the status of the building not being changed to unknown in response to the earthquake). When asked what adjustment to the robots model was required, 69% of all participants formulated an adjustment in behavioral terms (i.e., Inspect buildings in earthquake radius directly after an earthquake occurs). All other participants formulated an adjustment in terms of knowledge (i.e., increase level of priority of buildings in earthquake radius). After having been prompted with the Knowledge-rule LDP, the majority (81%) of the LDP group was still able to correctly formulate the required knowledge rule to improve the behavior of the robot, and expected the behavior to improve as a result of the rule (i.e., the robot will go to buildings hit by the earthquake). Thus, while most participants (69%) suggested behavioral adjustments to the robot's model rather than knowledge adjustments, when specifically asked for, all participants were able to formulate the knowledge required to improve the robot's model, mostly by using the terms learned during the Sitrep LDP.

During the Knowledge-rule LDP, participants were asked to specify a knowledge rule that, when implemented in the robot's knowledge model, would produce correct robot behavior after the earthquake. The participant received feedback on their knowledge-rule with an explanation based on the knowledge model (see chapter 4) on how this rule would affect its behavior. We asked participants how they evaluated the robot's explanation. In general, participants found the robot's explanation to be clear ($M = 4.20, SD = 0.68$), and indicated that it enabled them to better understand the behavior of the robot ($M = 4.20, SD = 0.86$). From the qualitative analysis (see Table 6) it becomes clear that people appreciated the explanation because of its visual, step-by-step presentation of how the robots behavior is established. Participants said it helped them to better understand how the robot will respond the next time an earthquake occurs.

8. Discussion

In this study we designed Learning Design Patterns (LDPs) for achieving co-learning in human-AI teams, and evaluated those in a wizard-of-Oz setting. Co-learning occurs from interactions that enable humans and AI agents to discover and learn about the task, themselves, and about each other. The LDPs that we designed describe and prescribe the learning of behaviors that are needed for handling difficult types of situations that often occur in a task. We designed two LDPs, consisting of interaction sequences that support co-learning of humans and AI, imposing several demands on both the human and AI partner. The effects of LDPs on learning and performance were evaluated within the context of a simulated urban-search-and-rescue-task. In particular, effects on the humans understanding of the AI partner, the human's perception of the human-AI collaboration, and the teams performance were investigated (see Section 8.1). We also investigated whether the method of predefining sequences of interactions in the form of LDPs is appropriate for designing effective learning for human-AI teams (see

Section 8.2).

8.1. Effects of LDPs on learning and team performance

Two LDPs were administered to the human-robot teams. The Sitrep LDP involved interactions intended to support the human developing an understanding why the robot showed erroneous task behavior in response to a particular event, thus threatening the teams performance. This LDP required the wizard-controlled robot to explain its behavior in intentional terms, and the human to respond to the explanation in order to increase understanding. The Knowledge-rule LDP involved interactions intended to support the human in teaching the robot how to act appropriately under such circumstances. This LDP required the human to create a rule based on concepts from the knowledge representation of the robot, while the robot has to incorporate this rule into its model and be able to feedback the behavioral consequences of the rule by means of an explanation. It was expected that engaging in these Learning Design Patterns would support team members to learn from each other, and to improve their performance in the second run (compared with a control condition in which teams did not engage in the LDPs). Although we did not find an effect on performance, we did find an effect on learning and understanding.

The interactions of the LDPs supported humans to better understand the robot. Participants were able to provide more accurate and detailed explanations of the robots behavior, when compared to participants that did not engage in the LDPs. Furthermore, the LDPs supported participants to develop a better awareness of the robots behavior, the teamwork, and the performance of the team (e.g., as shown by their opinions about the robots behavior, and the certainty of the observed robot behavior in the second run). It was expected that this better awareness would result in a more fluent collaboration between partners, and an improved team performance. This, however, was not found, which is in contrast to previous studies that found an association between awareness, understanding, and team performance (e.g., [Demir et al. \(2020\)](#); [Osofsky et al. \(2012\)](#)). One reason for our study not demonstrating this relationship may have to do with the relatively limited role of the LDPs on the overall performance on the USAR task. That is, although the LDPs addressed a piece of knowledge that is critical for developing awareness and understanding, this awareness and understanding contributed to the teams overall performance in a relatively minor manner. Thus, the effects of the LDPs may have been of too limited importance for the teams performance to demonstrate an effect.

Another reason for not finding an effect of LDPs on performance might be that participants experienced low team cohesion. Team cohesion is described as a bond that drives team members to remain motivated to work together to accomplish a set of goals ([Casey-Campbell and Martens, 2009](#)). In our task, participants could send commands to intervene with the behavior of the robot. However, there was little need for the participant to do so, as the robot performed all actions efficiently and autonomously. This could have elicited participants to assume a more supervisory control attitude, rather than a team member role. An indication for this is the relatively few communications initiated by the human.

8.2. Designing a learning context for human-AI teams

We have emphasized the need to study how human-AI teams jointly learn. Therefore, there is a need to develop tasks, environments, and procedures enabling such research. In this paper we have proposed such a context, which consists of a team task involving interdependencies between members, a dynamic knowledge representation of the AI partner to inform a realistic wizard-of-Oz protocol, and sequences of interactions in the form of Learning Design Patterns. As all learning is context dependent, we took efforts to embed LDPs in a context that is typical for human-AI teamwork, to enable their application for human-AI co-learning in other but similar contexts. The sequences of actions

and interactions in the LDPs were generically formulated. The Sitrep LDP is an intervention that encourages humans to look inside the 'brain' of their AI team partner, an interaction that intends to foster understanding of the team members thinking and reasoning. The Knowledge-rule LDP is an intervention that requires the human to think about what its AI-partner needs to learn, and to design and apply a knowledge intervention that fulfills that need. This interaction gives the human the opportunity to feed the AI agent with knowledge to be used when making its decisions. Such interventions should support collaboration and understanding over a longer period of time.

Further research in similar contexts is of course necessary to validate the generalizability of the suggested LDPs. Moreover, it should be taken into account that we have made several assumptions regarding the capabilities of an AI team member, e.g. on explainability and the presence of an explicit, ontology-based knowledge model. These assumptions should be considered when attempting to apply the LDPs in future research. Our propositions for LDPs present a start to explore the potential and value of the interventions for human-AI learning in context.

It is well known that effective team performance requires common ground between the team partners, for example in the form of a shared goal and a shared vocabulary. A formal knowledge representation was developed (see chapter 4) to provide our human-robot team with such a shared vocabulary. The LDPs in the study were based upon the concepts used in the knowledge representation. Studying the human-robot communication that was initiated by the LDPs enabled us to gain insight into the learning that took place in the human team member, and to assess whether shared understanding was achieved (i.e., whether the gained knowledge of the human aligned with that of the robot). Our work shows that a formal knowledge representation can be useful to facilitate communicative interactions in co-learning activities. We believe that such a human-understandable representation of knowledge should be used as basis for communication between human and AI team members, in order to enable them to establish common ground when working together.

As argued earlier, LDPs can only have a positive effect if the interactions address the needs for learning that are typical for tasks to be performed by human-AI teams. The need to view the AI agent as a true partner, the preparedness to improve as a team by joint collaboration, and the willingness to learn about the team partner are a few of them. Our implementation of a human-AI task environment did not always satisfy all these needs. A large part of the participants considered the robot to be a tool rather than a partner. The participants that engaged in the LDPs developed an understanding of the robot, and the USAR task required collaboration because of hard dependencies that were created between team members ([Johnson and Bradshaw, 2021](#); [Johnson et al., 2014b](#)). Still, the observation that participants viewed the effects of learning as a behavior change rather than as knowledge development, suggests that participants were not (yet) interested in the long term significance of their learning activities. Of course, participants were aware that their collaboration with the robot was limited to the duration of the experiment only. This may have had an influence on the participants attitude. In future research, it is important to design tasks in such a fashion that the human perceives the AI agent as a real team partner. This can for example be achieved by creating more soft interdependencies between the team members to encourage proactive helping, or by using psychological mechanisms such as described in [Nass et al. \(1996\)](#) to enhance team feeling. Moreover, learning should take place in a more natural way, meaning that learning interactions happen back and forth over a longer period of time. In our experiment learning was mostly done in a one-way isolated interaction, which is not how people learn in natural environments. Lastly, a limitation of this study is the use of a wizard-of-Oz technique to emulate the behavior of the robot. By using this method, we were unable to incorporate the human feedback (i.e., the knowledge rule) in the robot's knowledge and behavior model. Although we attempted to make the behavior of the AI partner as realistic as possible by modeling a knowledge base as well as a

goal hierarchy tree as basis for the wizard's protocol, the AI models remained static and thus did not support learning. Still, we do think the wizard-of-Oz approach provides a valuable first step to develop understanding of human experience and behavior in co-learning between human and AI. In a next study, it might be interesting to extend the research environment with a dynamic AI-model to study co-learning from both a human and AI perspective.

9. Conclusion

The rapid advancement of technology empowered by artificial intelligence is believed to bring forth new ecosystems in which human and AI act as complementing partners (Chui, 2017). 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). This co-learning may occur implicitly by experience during collaborations. It may also take place intentionally, by using Learning Design Patterns that elicit the interactions that produce learning in human-AI teams. In the present study, we designed two examples of LDPs: the Sitrep-LDP and Knowledge Rule-LDP, and implemented these in a human-AI co-learning testbed for research. The LDPs showed positive effects on human awareness and understanding of an AI agents behavior, but it may require additional efforts to advance improved awareness into better team performance. Based on experiences during our study, we identify several conditions for intentional co-learning to develop. First, it should be clear for the team why learning from each other is likely to benefit the teams functioning. That is, the context should provide an intrinsic motivation to learn. Second, the team should be supported in performing activities that provide opportunities to learn, such as after-action reviews in which team members exchange reflections and explanations of their behavior. Third, effective communication demands partners to use common concepts and a shared vocabulary. Humans tend to view behavior of themselves and of others in terms of everyday concepts such as beliefs, desires, and plans. This is often referred to as folk psychology (Horgan and Woodward, 1985). AI agents should therefore be equipped with a system that allows processing communication input from the human, and that enables outputting explanations in a form that can be understood by humans. Lastly, team members need to be confident that engaging in learning activities will help them to perform better. Being able to predict the effects of learning on team performance will support that conviction. When these conditions are incorporated into a co-learning research environment, they can provide the required opportunities to study intentional co-learning in human-AI teams.

CRedit authorship contribution statement

Tjeerd A.J. Schoonderwoerd: Conceptualization, Methodology, Formal analysis, Investigation, Writing – original draft, Writing – review & editing, Validation, Visualization. **Emma M. van Zoelen:** Conceptualization, Formal analysis, Investigation, Writing – original draft, Writing – review & editing, Validation, Visualization. **Karel van den Bosch:** Conceptualization, Writing – review & editing, Validation. **Mark A. Neerincx:** Validation, Writing – review & editing.

Declaration of Competing Interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

References

Ajoudani, A., Zanchettin, A.M., Ivaldi, S., Albu-Schäffer, A., Kosuge, K., Khatib, O., 2018. Progress and prospects of the human-robot collaboration. *Auton. Robots* 42 (5), 957–975.

Alexander, C., 1977. *A pattern language: Towns, buildings, construction*. Oxford university press.

van den Bosch, K., Schoonderwoerd, T., Blankendaal, R., Neerincx, M., 2019. Six challenges for human-AI co-learning. *International Conference on Human-Computer Interaction*. Springer, pp. 572–589.

Bradshaw, J.M., Sierhuis, M., Acquisti, A., Feltyov, P., Hoffman, R., Jeffers, R., Prescott, D., Suri, N., Uszok, A., Van Hoof, R., 2003. Adjustable autonomy and human-agent teamwork in practice: An interim report on space applications. *Agent autonomy*. Springer, pp. 243–280.

Broekens, J., Harbers, M., Hindriks, K., Van Den Bosch, K., Jonker, C., Meyer, J.-J., 2010. Do you get it? User-evaluated explainable BDI agents. *German Conference on Multiagent System Technologies*. Springer, pp. 28–39.

Buxbaum, H., Sumona, S., Kremer, L., 2019. An investigation into the implication of human-robot collaboration in the health care sector. *IFAC-PapersOnLine* 52 (19), 217–222.

Casey-Campbell, M., Martens, M.L., 2009. Sticking it all together: a critical assessment of the group cohesion–performance literature. *Int. J. Manag. Rev.* 11 (2), 223–246.

Chakraborti, T., Sreedharan, S., Zhang, Y., Kambhampati, S., 2017. Plan explanations as model reconciliation: moving beyond explanation as soliloquy. *arXiv preprint arXiv: 1701.08317*.

Cheah, C. W., 2008. Multi-site project management (MSPM)-an ontology supported methodology (on-line release).

Chui, M., 2017. *Artificial Intelligence the Next Digital Frontier*. McKinsey and Company Global Institute 47, 3–6.

De Graaf, M.M.A., Malle, B.F., 2017. How people explain action (and autonomous intelligent systems should too). 2017 AAAI Fall Symposium Series.

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. *Int. J. Hum. Comput. Stud.* 102436.

Dennett, D.C., 1989. *The Intentional Stance*. MIT Press.

van Diggelen, J., Johnson, M., 2019. Team design patterns. *Proceedings of the 7th International Conference on Human-Agent Interaction*, pp. 118–126.

Dillenbourg, P., Baker, M., Blaye, A., O'Malley, C., 1996. The evolution of research on collaborative learning. In: Spada, H., Reimann, P. (Eds.). *Learning in Humans and Machines*.

Feng, W., Zhuo, H.H., Kambhampati, S., 2018. Extracting action sequences from texts based on deep reinforcement learning. *arXiv preprint arXiv:1803.02632*.

Green, A., Huttenrauch, H., Eklundh, K.S., 2004. Applying the wizard-of-oz framework to cooperative service discovery and configuration. *RO-MAN 2004. 13th IEEE International Workshop on Robot and Human Interactive Communication (IEEE Catalog No. 04TH8759)*. IEEE, pp. 575–580.

Harbers, M., van den Bosch, K., Meyer, J.-J., 2010. Design and evaluation of explainable BDI agents. 2010 IEEE/WIC/ACM International Conference on Web Intelligence and Intelligent Agent Technology, Vol. 2. IEEE, pp. 125–132.

Harbers, M., Broekens, J., Van Den Bosch, K., Meyer, J.-J., 2010. Guidelines for developing explainable cognitive models. *Proceedings of ICCM*. Citeseer, pp. 85–90.

Harbers, M., Van Den Bosch, K., Meyer, J.-J., 2009. A methodology for developing self-explaining agents for virtual training. *International Workshop on Languages, Methodologies and Development Tools for Multi-Agent Systems*. Springer, pp. 168–182.

Hoffman, G., 2019. Evaluating fluency in human-robot collaboration. *IEEE Trans. Hum. Mach. Syst.* 49 (3), 209–218.

Holstein, K., Alevan, V., Rummel, N., 2020. A conceptual framework for human-AI hybrid adaptivity in education. *International Conference on Artificial Intelligence in Education*. Springer, pp. 240–254.

Horgan, T., Woodward, J., 1985. Folk psychology is here to stay. *Philos. Rev.* 94 (2), 197–226.

Johnson, M., Bradshaw, J.M., 2021. How Interdependence Explains the World of Teamwork. In: Lawless, W.F., Liinas, J., Sofge, D.A., Mittu, R. (Eds.), *Engineering Artificially Intelligent Systems: A Systems Engineering Approach to Realizing Synergistic Capabilities*. Springer International Publishing, Cham, pp. 122–146. https://doi.org/10.1007/978-3-030-89385-9_8.

Johnson, M., Bradshaw, J.M., Feltyov, P.J., Jonker, C.M., Van Riemsdijk, M.B., Sierhuis, M., 2014. Coactive design: designing support for interdependence in joint activity. *J. Hum.-Robot Interact.* 3 (1).

Johnson, M., Bradshaw, J.M., Feltyov, P.J., Jonker, C.M., Van Riemsdijk, M.B., Sierhuis, M., 2014. Coactive design: designing support for interdependence in joint activity. *J. Hum.-Robot Interact.* 3 (1), 43–69.

Klein, G., Woods, D.D., Bradshaw, J.M., Hoffman, R.R., Feltyov, P.J., 2004. Ten challenges for making automation a "team player" in joint human-agent activity. *IEEE Intell. Syst.* 19 (06), 91–95. <https://doi.org/10.1109/MIS.2004.74>.

Kruijff, G.-J.M., Janíček, M., Keshavdas, S., Laroche, B., Zender, H., Smets, N.J., Mioch, T., Neerincx, M.A., Diggelen, J.V., Colas, F., et al., 2014. Experience in system design for human-robot teaming in urban search and rescue. *Field and Service Robotics*. Springer, pp. 111–125.

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. *Proceedings of the Human Factors and Ergonomics Society Annual Meeting*. SAGE Publications Sage CA: Los Angeles, CA, pp. 1503–1507. Issue: 1.

Li, N., Matsuda, N., Cohen, W.W., Koedinger, K.R., 2015. Integrating representation learning and skill learning in a human-like intelligent agent. *Artif. Intell.* 219, 67–91. <https://doi.org/10.1016/j.artint.2014.11.002>.

Matheson, E., Minto, R., Zampieri, E.G.G., Faccio, M., Rosati, G., 2019. Human-robot collaboration in manufacturing applications: a review. *Robotics* 8 (4), 100.

- Miller, T., 2019. Explanation in artificial intelligence: insights from the social sciences. *Artif. Intell.* 267, 1–38.
- Mioch, T., Peeters, M.M.M., Neerincx, M.A., 2018. Improving adaptive human-robot cooperation through work agreements. 2018 27th IEEE International Symposium on Robot and Human Interactive Communication (RO-MAN). IEEE, pp. 1105–1110.
- Mitchell, T., Kisiel, B., Krishnamurthy, J., Lao, N., Mazaitis, K., Mohamed, T., Nakashole, N., Platanios, E., Ritter, A., Samadi, M., Settles, B., Cohen, W., Wang, R., Wijaya, D., Gupta, A., Chen, X., Saparov, A., Greaves, M., Welling, J., Hruschka, E., Talukdar, P., Yang, B., Betteridge, J., Carlson, A., Dalvi, B., Gardner, M., 2018. Never-ending learning. *Commun. ACM* 61 (5), 103–115. <https://doi.org/10.1145/3191513>.
- Nass, C., Fogg, B.J., Moon, Y., 1996. Can computers be teammates? *Int. J. Hum. Comput. Stud.* 45 (6), 669–678.
- Nikolaïdis, S., Hsu, D., Srinivasa, S., 2017. Human-robot mutual adaptation in collaborative tasks: models and experiments. *Int. J. Rob. Res.* 36 (5–7), 618–634.
- Niles, I., Pease, A., 2001. Towards a standard upper ontology. *Proceedings of the International Conference on Formal Ontology in Information Systems-Volume 2001*, pp. 2–9.
- Ososky, S., Schuster, D., Jentsch, F., Fiore, S., Shumaker, R., Lebiere, C., Kurup, U., Oh, J., Stentz, A., 2012. The importance of shared mental models and shared situation awareness for transforming robots from tools to teammates. *Unmanned Systems Technology XIV*, Vol. 8387. International Society for Optics and Photonics, p. 838710.
- Patterson, R.E., Pierce, B.J., Bell, H.H., Klein, G., 2010. Implicit learning, tacit knowledge, expertise development, and naturalistic decision making. *J. Cogn. Eng. Decis. Mak.* 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., 2020. Hybrid collective intelligence in a humanAI society. *AI & SOCIETY*. <https://doi.org/10.1007/s00146-020-01005-y>.
- Reber, P.J., Batterink, L.J., Thompson, K.R., Reuveni, B., 2019. Implicit learning: history and applications. *Implicit Learn.* 50, 16–37.
- Riek, L.D., 2012. Wizard of oz studies in hri: a systematic review and new reporting guidelines. *J. Hum.-Robot Interact.* 1 (1), 119–136. Steering Committee.
- Salas, E., Cannon-Bowers, J.A., Johnston, J.H., 1997. How can you turn a team of experts into an expert team?: Emerging training strategies. *Natural. Decis. Mak.* 1, 359–370.
- Sequeira, P., Alves-Oliveira, P., Ribeiro, T., Di Tullio, E., Petisca, S., Melo, F.S., Castellano, G., Paiva, A., 2016. Discovering social interaction strategies for robots from restricted-perception wizard-of-oz studies, Vol. 2016-April, pp. 197–204. <https://doi.org/10.1109/HRI.2016.7451752>.
- Sorensen, L.J., Stanton, N.A., 2016. Keeping it together: the role of transactional situation awareness in team performance. *Int. J. Ind. Ergon.* 53, 267–273.
- Sreedharan, S., Kambhampati, S., 2018. Handling model uncertainty and multiplicity in explanations via model reconciliation. *Proceedings of the International Conference on Automated Planning and Scheduling*, Vol. 28.Issue: 1.
- Stanton, N.A., Salmon, P.M., Walker, G.H., Salas, E., Hancock, P.A., 2017. State-of-science: situation awareness in individuals, teams and systems. *Ergonomics* 60 (4), 449–466. <https://doi.org/10.1080/00140139.2017.1278796>.
- Stout, R.J., Cannon-Bowers, J.A., Salas, E., 2017. The role of shared mental models in developing team situational awareness: implications for training. *Situational Awareness*. Routledge, pp. 287–318.
- Van Diggelen, J., Neerincx, M., Peeters, M., Schraagen, J.M., 2018. Developing effective and resilient human-agent teamwork using team design patterns. *IEEE Intell. Syst.* 34 (2), 15–24.
- Van Welie, M., Van Der Veer, G.C., Eliëns, A., 2001. Patterns as tools for user interface design. *Tools for Working with Guidelines*. Springer, pp. 313–324.
- Van Welie, M., Van der Veer, G.C., Eliëns, A., 1998. An ontology for task world models. *Design, Specification and Verification of Interactive Systems 98*. Springer, pp. 57–70.
- Wenskovitch, J., North, C., 2020. Interactive AI: designing for the two black boxes problem. *Computer* 53, 29–39.
- Woods, D.D., Tittle, J., Feil, M., Roesler, A., 2004. Envisioning human-robot coordination in future operations. *IEEE Trans. Syst. Man. Cybern. Part C* 34 (2), 210–218.
- Zimmerman, J., Forlizzi, J., Evenson, S., 2007. Research through design as a method for interaction design research in HCI. *Proceedings of the SIGCHI conference on Human factors in computing systems*, pp. 493–502.