



ARL-TR-10025 • Nov 2024



Intelligent Mission-Planning Technologies: Procedures and Findings from Phase 1 and 2 Experiments

**by Anthony L Baker, David Chhan, Bret L Kellihan,
Sarah Thomas, Aaron Necaie, Joshua Foldes, C Shawn Burke,
Katherine R Cox, and Joe T Rexwinkle**

DISTRIBUTION STATEMENT A. Approved for public release: distribution unlimited.

NOTICES

Disclaimers

The findings in this report are not to be construed as an official Department of the Army position unless so designated by other authorized documents.

Citation of manufacturer's or trade names does not constitute an official endorsement or approval of the use thereof.

Destroy this report when it is no longer needed. Do not return it to the originator.



Intelligent Mission-Planning Technologies: Procedures and Findings from Phase 1 and 2 Experiments

**Anthony L Baker, David Chhan, Katherine R Cox, and
Joe T Rexwinkle**
DEVCOM Army Research Laboratory

Bret L Kellihan
DCS Corporation

Sarah Thomas, Aaron Necaie, Joshua Foldes, and C Shawn Burke
University of Central Florida

REPORT DOCUMENTATION PAGE

1. REPORT DATE		2. REPORT TYPE		3. DATES COVERED	
November 2024		Technical Report		START DATE 9 October 2022	END DATE 29 September 2024
4. TITLE AND SUBTITLE Intelligent Mission-Planning Technologies: Procedures and Findings from Phase 1 and 2 Experiments					
5a. CONTRACT NUMBER W911NF-21-2-0279		5b. GRANT NUMBER		5c. PROGRAM ELEMENT NUMBER	
5d. PROJECT NUMBER ARL-22-156; ARL-24-025		5e. TASK NUMBER		5f. WORK UNIT NUMBER	
6. AUTHOR(S) Anthony L Baker, David Chhan, Bret L Kelliham, Sarah Thomas, Aaron Necaise, Joshua Foldes, C Shawn Burke, Katherine R Cox, and Joe T Rexwinkle					
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) DEVCOM Army Research Laboratory ATTN: FCDD-RLA-FD Aberdeen Proving Ground, MD 21005				8. PERFORMING ORGANIZATION REPORT NUMBER ARL-TR-10025	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)		10. SPONSOR/MONITOR'S ACRONYM(S)		11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT DISTRIBUTION STATEMENT A. Approved for public release: distribution unlimited.					
13. SUPPLEMENTARY NOTES ORCIDs: Anthony L Baker, 0000-0001-7163-4439; David Chhan, 0000-0003-2470-8663; Katherine R Cox, 0000-0002-8351-751X; Joe T Rexwinkle, 0000-0002-3856-101X					
14. ABSTRACT We are developing intelligent mission-planning technologies to predict performance and workload in mission areas using real-time data. These technologies will enhance mission planning and optimize Soldier capabilities during execution. Two studies, Phases 1 and 2, have been conducted to evaluate this capability. Phase 1 focused on combining participant feedback from after-action reviews with in-mission data to predict workload and performance in subsequent missions, validating the workload-prediction algorithm. Phase 2 expanded this concept to two participants working together and tested a new cooperative planning interface—demonstrating a learning pipeline that adapts in multi-agent scenarios. From these studies, we demonstrated that an evolving algorithm can generalize workload predictions across varying mission contexts. This approach enables the synthesis of multiple diverse data streams to accelerate the mission-planning process and optimize the performance of multi-human multi-agent teams. Ultimately, the technology aims to inform Soldiers of challenging situations in the battlespace, which will enhance mission outcomes in complex, dynamic future conflicts.					
15. SUBJECT TERMS human–autonomy teaming, mission planning, machine learning, ML, workload, software development, Humans in Complex Systems					
16. SECURITY CLASSIFICATION OF:				17. LIMITATION OF ABSTRACT	
a. REPORT UNCLASSIFIED	b. ABSTRACT UNCLASSIFIED	c. THIS PAGE UNCLASSIFIED	UU		61
19a. NAME OF RESPONSIBLE PERSON Anthony L Baker				19b. PHONE NUMBER (Include area code) (786) 427-5555	

STANDARD FORM 298 (REV. 5/2020)
Prescribed by ANSI Std. Z39.18

Contents

List of Figures	v
List of Tables	vii
1. Introduction	1
1.1 Summary	1
1.2 Background	2
2. Phase 1 Study	3
2.1 Overview and Goals	3
2.2 Methods	4
2.2.1 Tasks and Stimuli	4
2.2.2 Equipment and Setup	7
2.2.3 Questionnaires	9
2.2.4 Participants	10
2.2.5 Study Procedure	10
2.3 Results	13
2.3.1 Quantitative and Descriptives	13
2.3.2 Qualitative Data	18
2.3.3 Assessing the Learning Algorithm and Cost Predictions	19
2.4 Discussion	20
3. Phase 2 Study	21
3.1 Overview and Goals	21
3.2 Methods	22
3.2.1 Tasks and Stimuli	22
3.2.2 Equipment and Setup	26
3.2.3 Questionnaires	26
3.2.4 Participants	27
3.2.5 Study Procedure	28
3.3 Results	31
3.3.1 Quantitative and Descriptives	31

3.3.2	Qualitative Data	36
3.3.3	Assessing the Learning Algorithm and Cost Predictions	37
3.4	Discussion	39
4.	Conclusions and Next Steps	40
5.	References	42
	Appendix A. Cooperative Agreement Acknowledgment	43
	Appendix B. Questionnaires	45
	List of Symbols, Abbreviations, and Acronyms	51
	Distribution List	52

List of Figures

Fig. 1	An RCV in the simulated environment.....	5
Fig. 2	The left two targets should be marked friendly, as they are unarmed or holding noncombat objects. The right two targets should be marked hostile, as they are holding weapons (rifle and rocket-propelled grenade launcher, respectively).	5
Fig. 3	These three images depict how the cost map relates to the simulated environment. The left image is solely a portion of the environment shown from above. The rightmost image shows the cost map. The middle image shows the cost map overlaid onto the environment; this is a visual aid to understand how parts of the environment might relate to high or low cost. Cost maps are generated and shown at multiple points during the study.....	6
Fig. 4	These four images depict the terrain layers that factor into our cost map calculations. From left to right, the maps contain data for the terrain's foliage, urban areas, obstacles, and viewsheds. Because these have been previously determined for the entire simulation environment, we use these in workload calculations. We call these "precalculated terrain information."	7
Fig. 5	An approximate recreation of the setup used for the Phase 1 study. This was one of the two setups used for the Phase 2 study specifically, but the setups were functionally identical, so it is provided here as an example for the Phase 1 study as well.	8
Fig. 6	Sample of the Phase 1 UI. On the left is the target list and a small window where the user can view the thumbnails for targets requiring identification. The center displays the RCV's forward camera. The top right has a mini-map, and the bottom right contains buttons for switching between manual and autonomous navigation.....	9
Fig. 7	Mobility RT by scenario. This graphic is a box plot, displaying medians and IQRs for the factors, with swarm plots overlaid that indicate the overall data distributions for each factor. We see a slight difference between Scenarios 1 and 2.....	15
Fig. 8	Mobility travel time by scenario. Participants generally improved from one scenario to the next.....	16
Fig. 9	Target accuracy by scenario. There is a significantly notable upward trend for this metric.....	17
Fig. 10	Target RT by scenario. Participants performed most poorly on Scenario 1 but improved notably for subsequent scenarios.....	18
Fig. 11	Comparison of untrained, terrain-only cost map (left), post-Scenario 1 cost map (center), and post-Scenario 2 cost map (right). Differences between these cost maps indicate that the algorithm is indeed learning from each mission and changing its predictions based on participant	

	performance and environmental factors. The extent to which this is happening is subject to study in Phase 2 and beyond.	19
Fig. 12	Final cost maps for the simulated environment, resulting from training the model on the entire participant dataset. On the left, the cost map is represented for each individual pixel. On the right, a percentile filter is applied to “cluster” areas of similar workload, purely for visual representation of the data. Areas involving higher workload predictions appeared to relate most closely to obstacles and difficult terrain, such as the long river system on the left of the map, and the large urban area in the upper center of the map.....	20
Fig. 13	Participants sat next to each other and each had two computer monitors to use. Each monitor was associated with one RCV.	23
Fig. 14	In the planning phase, participants were newly able to cycle between the map view and the cost map view. The cost map color scale was also changed to accentuate differences between low-predicted workload areas (blues, with values close to 0.0), medium-predicted workload (grays with values around 0.5), and high-predicted workload areas (reds with values close to 1.0). This was instituted after reviewing feedback for the Phase 1 study.....	23
Fig. 15	Slider interface shown during the planning phase. On the left, the user has selected relatively high weights for the target and obstacle cost calculations, resulting in more red areas (high workload) in the cost map predictions. On the right, the user has selected lower weights for the target and obstacle calculations, leading to lower workload predictions (represented in blue). Not shown is the “generate route” button, which the user could press after adjusting the sliders. This would invoke the path-planning algorithm, optimizing the route for that vehicle based on the user’s preferred weights. Routes would change based on the weights selected, as the path-planning algorithm optimized for a low overall route cost.	24
Fig. 16	Example of a friendly vehicle, the Bradley Fighting Vehicle	25
Fig. 17	Sample of the Phase 2 UI. On the left is the target list and a small window that displays thumbnails for detected targets. The center displays the RCV’s forward camera. The right side contains a map of the area, and the user can toggle between a regular top-down map overlay or a cost map overlay. The bottom right contains the buttons for switching between manual and autonomous navigation.	26
Fig. 18	Explanation of Phase 2 study conditions	28
Fig. 19	Mobility RT box plot by scenario, displaying medians and IQRs for the factors. Overlaid swarm plots indicate the overall data distributions for each factor.	33
Fig. 20	Mobility travel time by scenario. Scenarios are generally similar-looking with largely overlapping plots, suggesting no improvement was seen in this factor across scenarios.	34

Fig. 21	Target accuracy by scenario. Scenarios are generally similar-looking with largely overlapping plots, suggesting no improvement was seen in this factor across scenarios.....	35
Fig. 22	Target RT by scenario. A downward trend is visually evident and supported by statistical tests, suggesting participants got faster at reacting to targets over the course of several missions. Occasional outliers result from participants sometimes forgetting to respond to a target for several minutes, then coming back to it later. The outliers in this plot were left in as they were found to be genuine mistakes rather than data errors, though some more distant outliers were removed from the plot and analysis.....	36
Fig. 23	Comparison of untrained, terrain-only cost map (left) with cost maps generated using data from every subsequent scenario. Differences between these cost maps indicate that the algorithm is indeed learning from each mission and changing its predictions based on participant performance and environmental factors, even as Phase 2 moved the learning pipeline into the multi-human multi-agent context.....	38

List of Tables

Table 1	Procedure and timeline for Phase 1 study	13
Table 2	Descriptive statistics, aggregated across all scenarios.....	14
Table 3	Descriptive statistics by scenario. Values indicate mean (SD). For each metric, asterisks and carats indicate significant differences on a Kruskal–Wallis test at $p < 0.05$ between cells containing the same symbol.....	14
Table 4	Procedure and approximate timeline for Phase 2 study	31
Table 5	Descriptive statistics for on and off conditions.....	32
Table 6	Descriptive statistics by scenario. Table values indicate mean (SD)..	32
Table 7	Total and average path costs for all scenarios	39

1. Introduction

1.1 Summary

U.S. Army Combat Capabilities Development Command (DEVCOM) Army Research Laboratory (ARL) researchers, in conjunction with partners at the University of Central Florida, are developing an intelligent mission-planning technology that can predict performance and workload across a mission area using models trained on opportunistically sensed in-mission information and mission-relevant environmental features. This capability will provide key information to improve mission planning and enable optimization of the moment-to-moment demands on the Soldier during mission execution. We have thus far conducted two research studies, called Phases 1 and 2, developing and evaluating this capability. The purposes of this report are to detail the procedures and findings of both studies and conclude with a discussion of the planned follow-up research.

The key objective of the Phase 1 study, which concluded in FY23, was to determine how participant feedback in an after-action review (AAR) could be combined with in-mission data to predict participant workload and performance in consecutive missions. This was a crucial first step to validate the closed-loop design of the workload-prediction algorithm, which processes mission and AAR data to inform subsequent mission planning, which then leads to new mission and AAR data, and so on. The key objectives of Phase 2 in FY24 were to expand the concept to two participants working together to plan and execute missions using the intelligent mission-planning technology and test a new cooperative mission-planning user interface (UI). This phase proved to be the concept of a learning pipeline that adapts mission-to-mission during multi-human, multi-agent scenarios and lays the groundwork for further expansion and refinement by evaluating the extent to which the planning technologies improve performance outcomes in multi-human, multi-agent teams.

The primary contributions of this work are demonstrating that an algorithm that evolves mission-to-mission can be used to predict generalized workload across changing mission areas and that this learning pipeline works in the context of a multi-human, multi-agent simulation. The complexities of the future operating environment will require effective synthesis of disparate data streams to inform forward projections of decisions and actions, and this initial work has demonstrated an approach to doing so. In the ideal use case, the intelligent mission-planning technology will be able to inform Soldiers of areas in the battlespace that will involve challenging situations so Warfighters can prepare and account for those during mission planning. Future studies will take necessary steps toward this goal.

1.2 Background

To meet the evolving challenges of the future operating environment, the U.S. Army is developing concepts that will see more widespread implementation of intelligent technologies to support Warfighter activities. A key component of the military's modernization efforts is understanding how to better facilitate human–autonomy teaming so human operators and autonomous systems can collaborate effectively. To support optimal performance of human–autonomy teams, it will be crucial to provide Soldiers with high-quality information during the mission-planning process. The larger goal of this research effort is to develop an intelligent mission-planning technology that can synthesize terrain information and previous mission data to predict and generalize performance characteristics across an entire battlespace.

Prior work involving human–agent teams in a cooperative control task demonstrated that when agents can adapt to human behaviors, teams are able to achieve a high level of performance more quickly than if the agents are nonadaptive and have a static behavior policy (Li et al. 2021). In that study, agents trained on data generated from human-based self-play were able to adapt to human teammate behaviors, adopting optimal policies for varied scenarios and allowing the teams to quickly reach a high degree of proficiency. In contrast, the performance of the teams in the non-adaptive conditions took much longer to reach an equivalent level of performance on the task, suggesting that adaptive autonomous capabilities can significantly augment human–autonomy teaming performance.

AARs consist of a structured approach to providing feedback on training or combat missions. Their goals are identifying and correcting deficiencies, augmenting strengths, and focusing on the performance of specific training objectives once a mission is complete (Brewer et al. 2019). AARs are targeted at debriefing team members to aid learning from prior performance (Smith-Jentsch et al. 1998), allowing teammates to practice resolving conflicts in a productive manner (Marks et al. 2001; De Dreu and Weingart 2003), and enabling the development of accurate mental models of problems, tasks, spaces, and shared knowledge among team members (Endsley 1995; Blickensderfer et al. 1997). The advantages of AARs warrant further study of how they can be used to augment human–autonomy team performance across multiple missions.

It will also be critical for autonomous Army vehicles to be able to determine the traversability of their environment and help compute optimal routes toward their goals (Wulfmeier et al. 2017). Route planning must involve balancing many factors such as terrain features, reported locations of hostile forces, availability of cover, and so on. Thus, there is a clear need for mission-planning tools that can effectively

learn from previous missions and Soldier inputs to help develop optimal plans for future scenarios.

Altogether, an intelligent mission-planning technology promises to enhance the processes involved with synthesizing disparate data streams, giving Army leaders optimal information and efficient summaries of complicated situations and accelerating the rate at which the Army can adapt to the complex, multidomain battlespace of the future.

2. Phase 1 Study

2.1 Overview and Goals

Our research on intelligent mission-support technology began with a 2019 project that had two phases. In the first phase, participants completed two tasks: they navigated vehicles through a simulated environment and identified targets. Target identification was assisted by simulated aided target recognition (AiTR) which had predefined strengths and weaknesses at different areas in the environment. The second phase had participants conduct AARs during which they stepped through the recently completed exercise in short intervals and indicated which camera feeds on the vehicles they were operating were most useful for each task. The data were then used to develop a cost map that indicated areas with expected high or low workload based on 1) the number of camera feeds that were likely important at any given point on the environment and 2) the extent to which the AiTR was able to assist. The simulated environment used in the ARL laboratory has been changed and upgraded since 2019, and the human-agent teaming AAR process has similarly been revamped to be significantly faster than the one used in 2019. Therefore, the present study plans to replicate and build upon the 2019 work (which was unpublished) and extend it using the new simulation environment and AAR procedure.

Phase 1 of this research line was planned to be a smaller study that would allow us to validate our methods and algorithms to eventually progress toward multi-human, multi-agent teams working with the intelligent planning technologies. The key objectives of Phase 1 were as follows:

- 1) To replicate and extend the 2019 research, with an updated simulation environment and AAR procedure.
- 2) To demonstrate the proof-of-concept for intelligent mission-planning tools (i.e., algorithm-generated workload predictions based on prior mission data and AAR feedback), and to gather initial, nonexperimental insights into

how in-mission performance and workload might change as mission-planning tools are updated.

Phase 1 was not envisioned with an experimental design, as the intelligent planning tools first needed to be proved and demonstrated before eventually assessing their efficacy in future studies. Thus, this study did not contain separate control and experimental conditions. Instead, all participants would complete all scenarios under the same circumstances to observe whether the learning pipeline and mission-to-mission structure of the overall study was functional. This would allow us to gather initial insights into how the mission-planning process and learning pipeline worked together.

2.2 Methods

This section details the overall structure of the study, discussing the tasks, equipment, questionnaires, participant pool, and procedure.

2.2.1 Tasks and Stimuli

Participants acted as the operator of two robotic combat vehicles (RCVs; Fig. 1) that are driving through an area while also identifying potential targets as friendly or hostile. For the driving task, the RCVs had autonomous mobility that allowed them to travel predetermined paths. Along the way, vehicles encountered scripted obstacles, terrain, or similar mobility threats that required participants to take over manual control and drive the RCVs back onto the planned route where they re-engaged the autonomous driving. Therefore, participants' goals in the driving task were to monitor autonomous driving in both RCVs and take over driving as quickly as possible when an obstacle was encountered so the RCV is never a sitting target for enemy forces. For Phase 1, mobility obstacles were preset at certain areas in each scenario to ensure some consistency between participants' experiment runs.



Fig. 1 An RCV in the simulated environment

For the target recognition task, the RCVs had a simulated AiTR system that detected potential targets in the environment (Fig. 2) and provided a classification of each as friendly or hostile. In this task, the participant's goal was to review the classification of each target by the AiTR and either accept or revise the system's classification. Some targets were more difficult to correctly identify than others (e.g., by being elevated or partially obscured). Targets were preset in scattered areas around the environment.



Fig. 2 The left two targets should be marked friendly, as they are unarmed or holding noncombat objects. The right two targets should be marked hostile, as they are holding weapons (rifle and rocket-propelled grenade launcher, respectively).

The autonomous navigation system operated the vehicles through predetermined routes in each scenario. In Scenarios 1 and 2, the routes for both RCVs were preset.

In Scenario 3, participants were given three route options for each RCV and were allowed to choose one for the upcoming mission. The route options prioritized different things, such as speed, safety, or more visibility of targets. Cost maps, which helped participants to visually understand where certain parts of the mission may incur a high or low workload, were used to aid planning in Scenarios 2 and 3 and route selection in Scenario 3. The cost maps are described in more detail later in this section.

After each mission, an AAR was conducted to review the previous mission. The AAR considered the preceding mission and allowed the participant to rate their subjective workload experienced at the moment they identified each target to indicate how difficult it was to correctly identify it (Appendix B). They did the same for each time they encountered a mobility obstacle. This information, combined with performance of the previous mission and local terrain data near user tasks, was fed into an ML algorithm that then classified areas of the map as high or low workload areas, and these data were shown in a cost map. This cost map was rendered as a top-down map of the next mission area, with certain areas of the map shaded in different colors indicating the predicted cognitive workload (Fig. 3). Terrain data is separated into four layers that factor into the cost map calculations (Fig. 4): foliage, urban areas, obstacles (impassable objects such as fences, walls, etc.), and viewsheds (i.e., the extent to which given areas have line-of-sight to other areas).

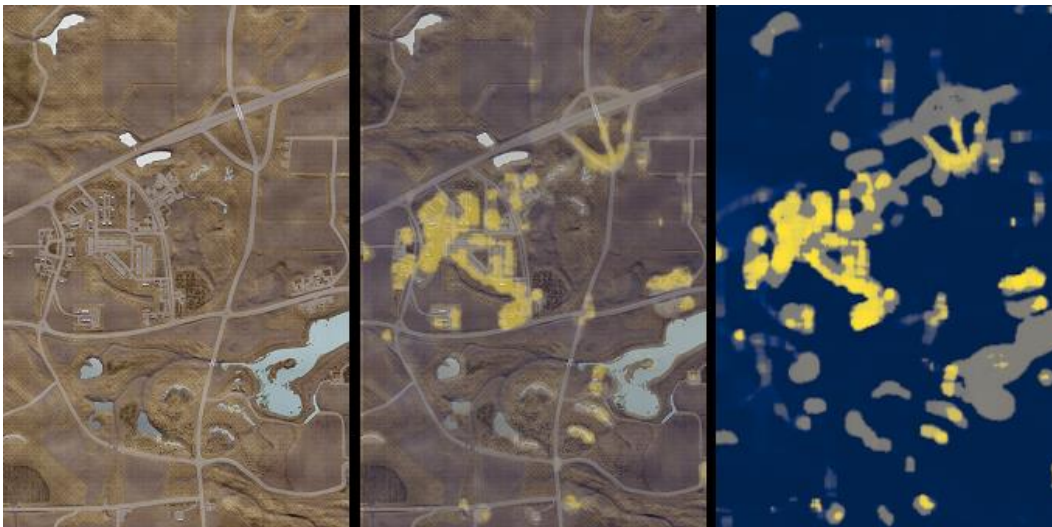


Fig. 3 These three images depict how the cost map relates to the simulated environment. The left image is solely a portion of the environment shown from above. The rightmost image shows the cost map. The middle image shows the cost map overlaid onto the environment; this is a visual aid to understand how parts of the environment might relate to high or low cost. Cost maps are generated and shown at multiple points during the study.

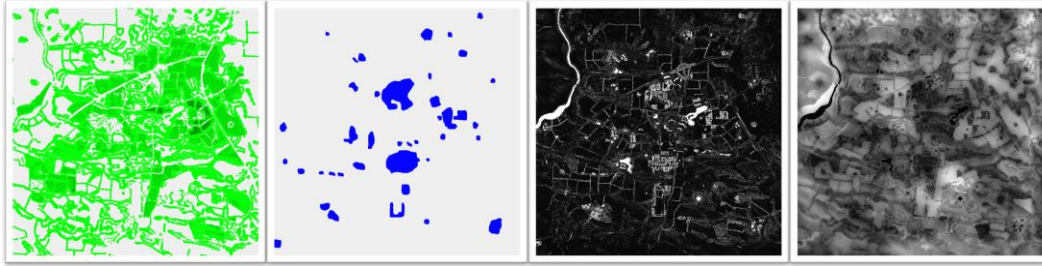


Fig. 4 These four images depict the terrain layers that factor into our cost map calculations. From left to right, the maps contain data for the terrain’s foliage, urban areas, obstacles, and viewsheds. Because these have been previously determined for the entire simulation environment, we use these in workload calculations. We call these “precalculated terrain information.”

The purpose of the cost map was to allow participants to mission-plan using tools that showed them areas in which they may experience more or less workload. Before any missions begin, an initial cost map is generated based only on preset predicted workload and performance values assigned to terrain data (Fig. 4), and as the participant completes each mission and AAR, their performance and reported workload are used to update the ML algorithm that populates the cost map. The algorithm considers terrain data (e.g., foliage, urban areas), performance in recovering from mobility failures and identifying targets, and participant workload data collected from the AAR. All the factors related to terrain, mobility failures, and AiTR data can be calculated a priori for any area of the simulated environment because we have those ground truth data from developing the scenarios.

All stimuli were created using Unreal Engine. All scenarios took place in a simulated environment that was used in prior ARL studies. Targets for identification (e.g., simulated people) were assets that have been used either in prior ARL studies or were otherwise developed for the simulation.

2.2.2 Equipment and Setup

The equipment required for the study consisted of a high-performance desktop computer, two 15-inch touchscreen monitors mounted next to each other, a mouse and keyboard, and a Logitech G920 steering wheel and pedal set for driving the RCV (Fig. 5). Each monitor rendered the interface for one RCV (Fig. 6).



Fig. 5 An approximate recreation of the setup used for the Phase 1 study. This was one of the two setups used for the Phase 2 study specifically, but the setups were functionally identical, so it is provided here as an example for the Phase 1 study as well.



Fig. 6 Sample of the Phase 1 UI. On the left is the target list and a small window where the user can view the thumbnails for targets requiring identification. The center displays the RCV's forward camera. The top right has a mini-map, and the bottom right contains buttons for switching between manual and autonomous navigation.

2.2.3 Questionnaires

Participants completed four questionnaires at various points throughout the study: the NASA Task Load Index (NASA-TLX) Short, a custom workload questionnaire, a driving experience questionnaire, and a simulation design feedback questionnaire.

The NASA-TLX Short is a questionnaire in which people self-report their perceived workload demands on six subscales: mental demand, physical demand, temporal demand, performance, effort, and frustration (Hart and Staveland 1988). The short version of the NASA-TLX only consists of the first section, where six items are used to gather an unweighted score for each subscale. For each subscale factor, participants respond on a scale ranging from 0 (low) to 100 (high) in five-point increments, and all the subscales are averaged to make up an overall workload score (see Appendix B).

The workload questionnaire was given after each mission stage to determine participants' level of workload difficulty they experienced while identifying each target or dealing with each mobility obstacle (Appendix B). The questions are on a seven-point scale, from the lowest possible workload (1) to the highest (7).

The driving experience questionnaire was given after the informed consent process. It asks three questions intended to gather information about participants' driving experience. This was aimed to help provide context for their interactions with the

navigation and mobility tasks to better understand performance variability (see Appendix B).

The simulation design feedback questionnaire was given to participants after they completed the final AAR of the study. It was intended to gather feedback from participants that aided in the development of the Phase 2 research study (discussed later in this report). See Appendix B for the questionnaire.

2.2.4 Participants

Thirty-six participants were recruited from the general population around the University of Central Florida area and completed the study. We primarily targeted the student body available on the Sona Systems participant recruitment and management system. The eligibility criteria for participants were as follows. Participants must

- be aged 18–45 years,
- have normal or corrected-to-normal vision required,
- not have any color vision difficulties,
- be able to view and interact with first-person video games (e.g., driving, shooting) without experiencing motion sickness, and
- not have photosensitive epilepsy.

Participants were compensated \$15/h of participation, rounded up to the nearest hour, up to a maximum of \$60 for 3.5 h. Participants were paid via Amazon gift cards in \$15 increments.

2.2.5 Study Procedure

2.2.5.1 Consent and Training

Participants arrived on site and read the informed consent document. Next, they filled out the driving experience questionnaire. Participants then completed a training session where they were introduced to the simulator, controls, and expected tasks. They were familiarized with the layout of the screen and how to verify or modify the AiTR's classification of targets as hostile or nonhostile. They were also instructed on how to drive the RCVs with manual controls and change between autonomous and manual modes.

The main experimental session consisted of three scenarios, each with a planning stage, a mission stage, and an AAR stage. Participants were tasked with monitoring

and correcting autonomous systems as they identified targets and navigated through a simulated environment.

2.2.5.2 Scenario 1

The Scenario 1 planning stage began with participants being shown the map of the environment with routes drawn in the upcoming mission area. Participants were briefed on their goal of monitoring and supporting autonomy in identifying targets and navigating the RCVs. The map of the mission area showed participants the routes that each RCV would take. The participant was also able to view the cost map display to see areas of predicted low or high workload. In Scenario 1, the cost map was based only on preset workload and performance estimates generated by the research team for different terrain features. For example, very steep elevation or ditches might be bright yellow because they are likely to cause mobility failures (and thus high workload), whereas cooler colors were representative of areas of more open or flat terrain. Participants were told that the map indicates the level of workload expected based on terrain data, such that more obstacles may lead to mobility failures or areas of high urban density may make target detection more difficult.

The Scenario 1 mission stage consisted of participants monitoring both RCVs as they navigated along the predetermined routes. Targets routinely became visible and populated the AiTR feed, and participants had to accept the AiTR's classification or revise them based on the image thumbnail and what they were able to see through the RCV camera feed. Multiple times (~4 to 5) throughout the mission, an RCV encountered a mobility obstacle and the participant had to take manual control and drive the vehicle out of the obstacle and back onto the mission route. As stated earlier, the locations of obstacles were predetermined along each route during the scenario development process.

The Scenario 1 AAR stage began after the participant completed the mission. Participants completed the NASA-TLX questionnaire to indicate their overall workload during the Scenario 1 mission. Next, participants reviewed each target that their vehicle encountered. An image was provided of the target to aid their recall. The participant reported their perception of the workload required to identify the target and the responses were captured by the system to feed the cost-map algorithm. The mobility obstacles were reviewed in a similar manner.

After the Scenario 1 AAR stage, the ML algorithm was trained on data from the mission and AAR to update the cost map for mission planning while the next scenario was prepared by the experimenter. During this time, the participant was offered a short break.

2.2.5.3 Scenario 2

Scenario 2 occurred in a different area of the map. The Scenario 2 planning stage consisted of the same format as in the Scenario 1 planning stage; however, the cost map was updated based on information from the previous mission and AAR to make workload predictions for the environment. While participants were able to view the entire map, they were shown the specific routes that the two RCVs would take in Scenario 2.

The Scenario 2 mission and AAR stages proceeded in the same manner as those in Scenario 1.

2.2.5.4 Scenario 3

Scenario 3 occurred in a different part of the map from the first two scenarios. The Scenario 3 planning stage had an updated cost map and the participant was given three route options for each RCV. The routes each had a different priority, such that each route incurred a higher predicted degree of workload from one source (mobility obstacles, AiTR targets, or preset values related to terrain/environment features). In other words, one route had a higher predicted degree of workload stemming from the mobility obstacles that were encountered; one had higher predicted workload stemming from the features of potential targets (high elevation, obscured behind cover, etc.) that were encountered along the route; and one had higher predicted workload stemming from the terrain or environment features of the route (e.g., the route went through more difficult terrain). Participants chose the route for each RCV based on their understanding of the cost map, and then the mission began.

The Scenario 3 mission proceeded in the same manner as those before, and a Scenario 3 AAR was conducted to collect additional data that helped algorithm development. Following the AAR, the participant was asked to provide feedback on the study paradigm that would aid development of the larger Phase 2 follow-up study. The study was expected to last approximately 3 h, but occasionally ran as short as 2.25 h and as long as 3.5 h (the maximum allowed duration). Table 1 provides the procedure and timeline for the Phase 1 study.

Table 1 Procedure and timeline for Phase 1 study

Event	Event time (min)	Cumulative time (h and min)
Informed consent and questionnaire	10	0 h 10 min
Training	10	0 h 20 min
Scenario 1 planning	10	0 h 30 min
Scenario 1 mission	~10	0 h 40 min
Scenario 1 AAR and questionnaires	~25	1 h 5 min
<i>Break, Scenario 2 setup</i>	5	1 h 10 min
Scenario 2 planning	10	1 h 20 min
Scenario 2 mission	~10	1 h 30 min
Scenario 2 AAR and questionnaires	~25	1 h 55 min
<i>Break, Scenario 3 setup</i>	5	2 h 0 min
Scenario 3 planning	10	2 h 10 min
Scenario 3 mission	~10	2 h 20 min
Scenario 3 AAR and questionnaires	~30	2 h 50 min
<i>Debrief</i>	5	2 h 55 min

2.3 Results

A series of analyses of variance were conducted using the three missions as three levels of the independent variable and mission performance factors (number of targets identified, speed of identifying targets, reaction time before taking over manual control) as dependent variables. Qualitative analyses were performed on the feedback questionnaire and observations from the researchers conducting the study also fed insights that were used to expand on the Phase 1 paradigm. We also evaluated the algorithm's cost-map generation between scenarios. Sections 2.3.1 through 2.3.3 detail each of these analyses.

2.3.1 Quantitative and Descriptives

We quantified mission performance based on several factors: 1) the accuracy of the user's target labels, 2) the response time (RT) to the target labels, 3) the time required to navigate RCVs out of mobility obstacles, and 4) the RT to mobility obstacles.

Descriptive statistics for the performance metrics—mobility RT, mobility travel time, target accuracy, and target RT—are reported in Tables 2 and 3.

Table 2 Descriptive statistics, aggregated across all scenarios

Metric	Mean	SD
Mobility RT (s)	7.46	11.09
Mobility travel duration (s)	30.20	16.71
Target accuracy (%)	61.69	7.02
Target RT (s)	12.06	11.00

Table 3 Descriptive statistics by scenario. Values indicate mean (SD). For each metric, asterisks and carats indicate significant differences on a Kruskal–Wallis test at $p < 0.05$ between cells containing the same symbol.

Scenario	Mobility RT (s)	Mobility travel time (s)	Target accuracy (%)	Target RT (s)
1	9.84 (16.44)*	37.79 (20.01)*^	56.65 (4.95)*	18.09 (17.42)*^
2	5.34 (5.23)*	26.71 (13.40)*	58.85 (3.43)^	8.75 (1.88)*
3	7.26 (8.47)	26.08 (13.79)^	70.00 (3.38)*^	9.36 (2.92)^

A few outliers beyond three standard deviations (SDs) were excluded from the dataset. Assessments of normality indicated the data did not fit normal distributions, and Levene’s tests were conducted on the study factors (mobility RT, mobility travel time, target accuracy, target RT) to assess the assumption of equal variances across scenarios. Most Levene’s tests were significant. Given the nonparametric data and the presence of unequal variances for some factors, we opted for Kruskal–Wallis H tests to assess whether there were significant differences in the medians of each factor across the three scenarios.

For Mobility RT (Fig. 7), the Kruskal–Wallis test revealed a significant difference across scenarios ($H = 6.96$, $p = 0.0308$), and pairwise comparisons found significant differences between Scenarios 1 and 2, but no other scenarios produced a notable difference.

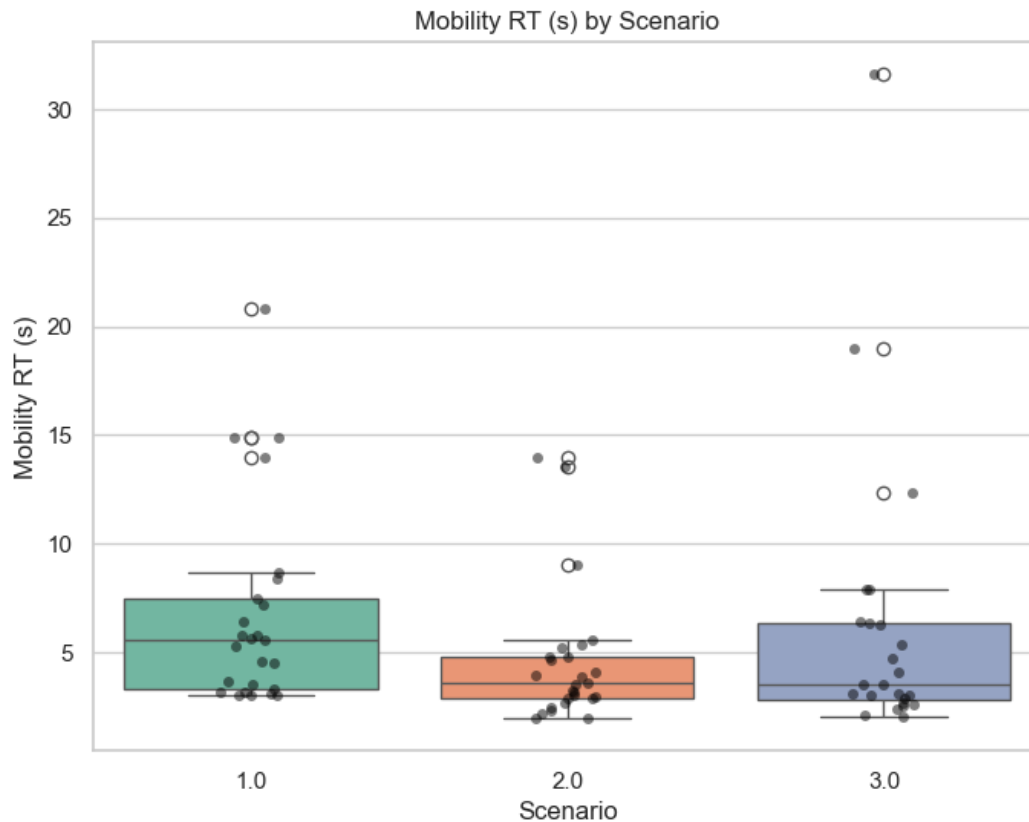


Fig. 7 Mobility RT by scenario. This graphic is a box plot, displaying medians and IQRs for the factors, with swarm plots overlaid that indicate the overall data distributions for each factor. We see a slight difference between Scenarios 1 and 2.

For mobility travel time (Fig. 8), more significant differences were observed ($H = 16.39$, $p < 0.001$), this time between Scenarios 1 and 2 ($p = 0.0048$) and Scenarios 1 and 3 ($p = 0.0005$). Generally, this seems to indicate that participants were quite slow in Scenario 1 and got better in Scenarios 2 and 3. This may have resulted from the Phase 1 setup where participants did not have a training scenario to complete. As a result, the Phase 2 study added a training scenario.

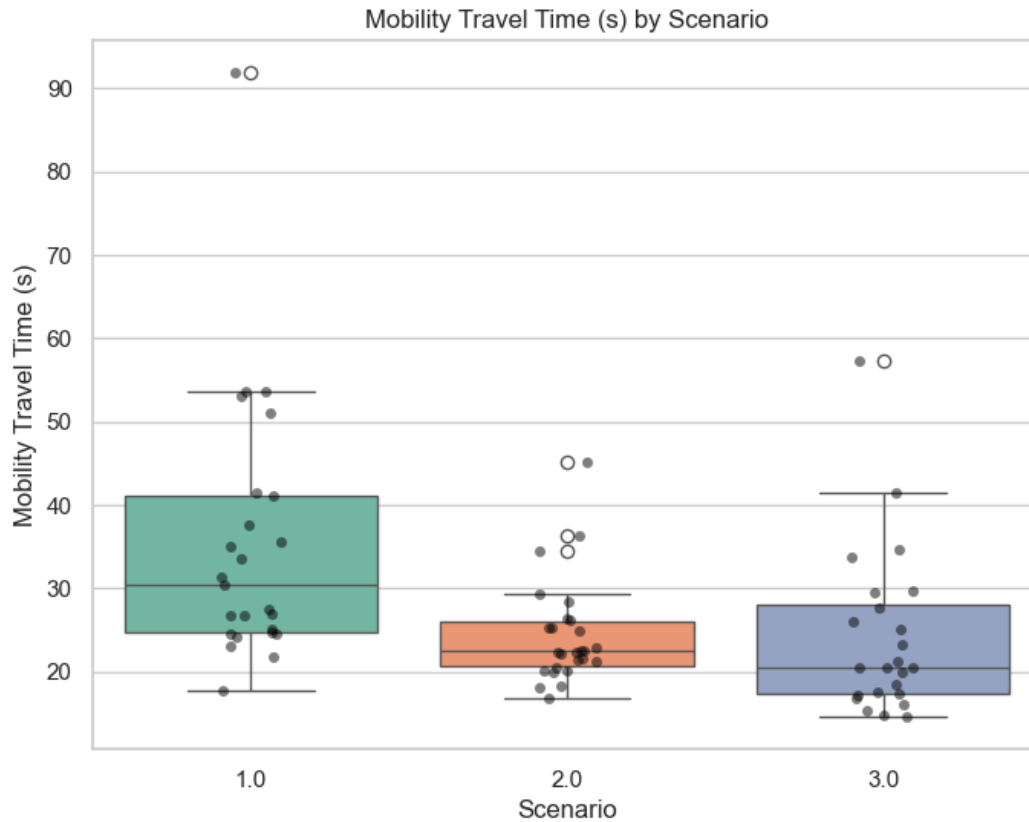


Fig. 8 Mobility travel time by scenario. Participants generally improved from one scenario to the next.

For target accuracy (Fig. 9), analysis indicated significant differences between scenarios ($H = 49.61, p < 0.001$). Posthoc tests found differences between Scenarios 1 and 3 ($p < 0.001$) and between Scenarios 2 and 3 ($p < 0.001$). There was a significant upward trend in participants' target accuracy by scenario.

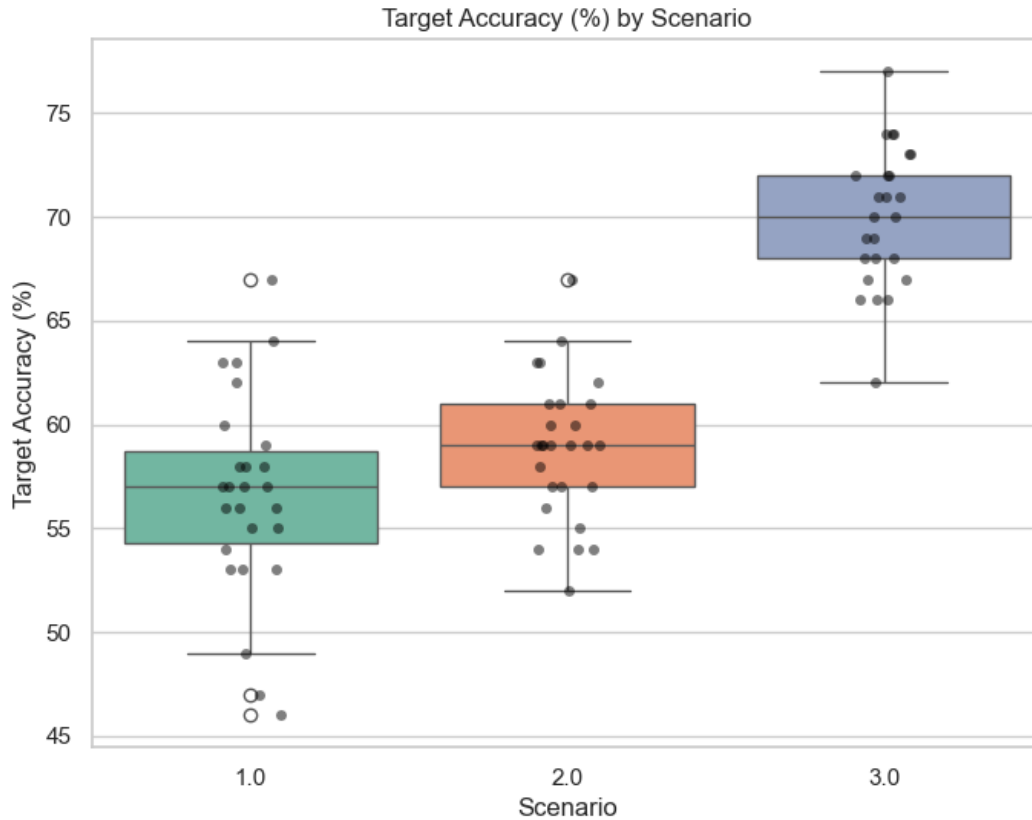


Fig. 9 Target accuracy by scenario. There is a significantly notable upward trend for this metric.

For the last analysis, we also found significant differences between scenarios on target RT ($H = 28.83$, $p < 0.001$). Significant differences were found between Scenarios 1 and 2 ($p < 0.001$), and between Scenarios 1 and 3 ($p < 0.001$), again indicating a trend of improvement across scenarios (Fig. 10).

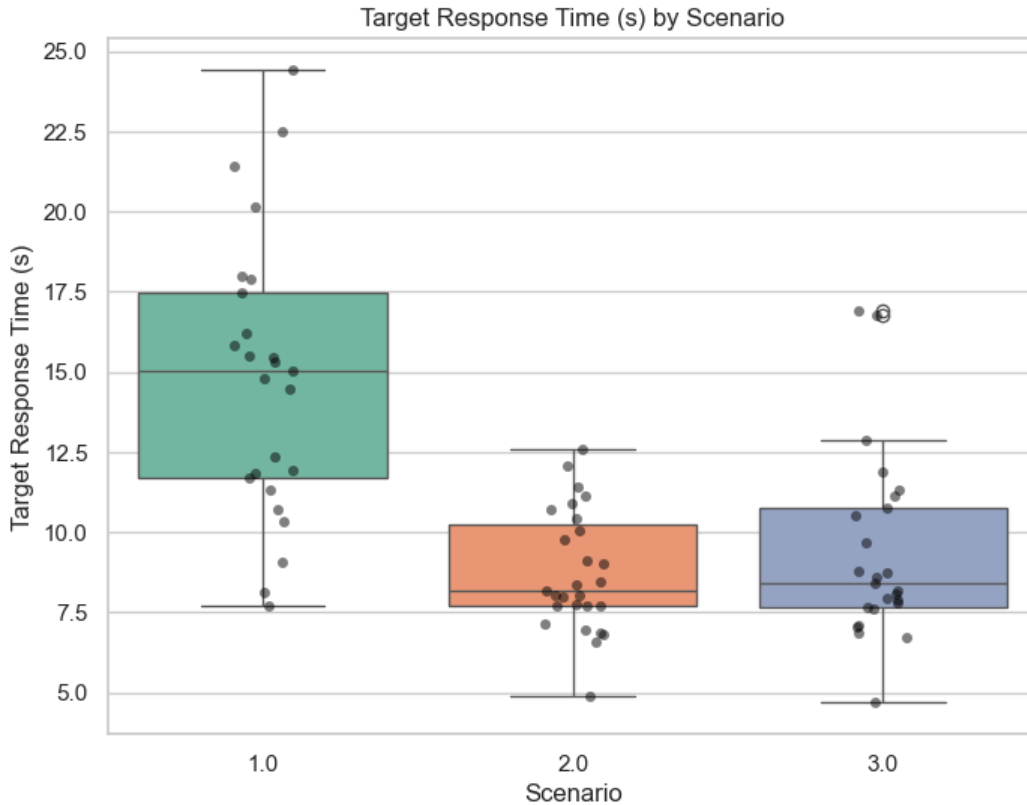


Fig. 10 Target RT by scenario. Participants performed most poorly on Scenario 1 but improved notably for subsequent scenarios.

The analyses overall suggest that participants generally improved on each of these metrics as they completed more scenarios, indicating a general practice effect as they become more accustomed to the tasks. Phase 1 did not include a separate scenario solely for training purposes, so we added that to the design of the Phase 2 study to reduce some of the practice effects observed across these Phase 1 analyses.

2.3.2 Qualitative Data

Qualitative analysis of feedback provided by participants after the study provided further insight into the utility of the intelligent mission-planning technology and the scenarios. 75% of participants found the cost map “somewhat useful” or better, but 25% rated it “not very useful.” Regarding specific features, 71% of participants called out the color gradients of the cost map and the ability to see the RCVs’ planned routes superimposed over the cost map as being very helpful to planning. However, 20% of participants wished to see the cost map during the mission phase (and not just during the planning phase). Besides that, we received 18 suggestions on minor adjustments to the interface that would improve the planning process, such as making a distinction between costs associated with the target identification task versus costs associated with mobility, so that operators can better plan around

the demands stemming from the two types of tasks. These insights were carried into the development process for the Phase 2 study.

2.3.3 Assessing the Learning Algorithm and Cost Predictions

We next reviewed the ability of the cost map to evolve mission-to-mission. For the first cost map shown to participants (i.e., before any missions have been completed and any mission data is available for analysis), only precalculated terrain information is available (see Fig. 4); thus, they are the only thing that affect the algorithm's workload predictions. We then trained two more cost maps on the entire participant datasets for the first and second Scenarios, resulting in a total of three cost maps that could be reviewed for cross-mission changes. For this review, we left out a post-Scenario 3 cost map, because participants could select significantly different routes that affected the obstacles and target placements. Given the small sample size in Phase 1, we found it likely that the post-Scenario 3 dataset would not have sufficient data to effectively train a cost map or be reliably comparable to the cost maps generated for the first two missions when all participants underwent the same routes.

Figure 11 shows a comparison of the Baseline cost map, the post-Scenario 1 cost map, and the post-Scenario 2 cost map. There is a visible difference between the workload predictions across missions, suggesting that some factors relating to the changing missions, or participants' performance in them, is indeed spurring the algorithm to update its predictions as missions are completed. This suggests that the algorithm and cross-mission learning pipeline are indeed achieving our goals of providing updated predictions across situations and environments.

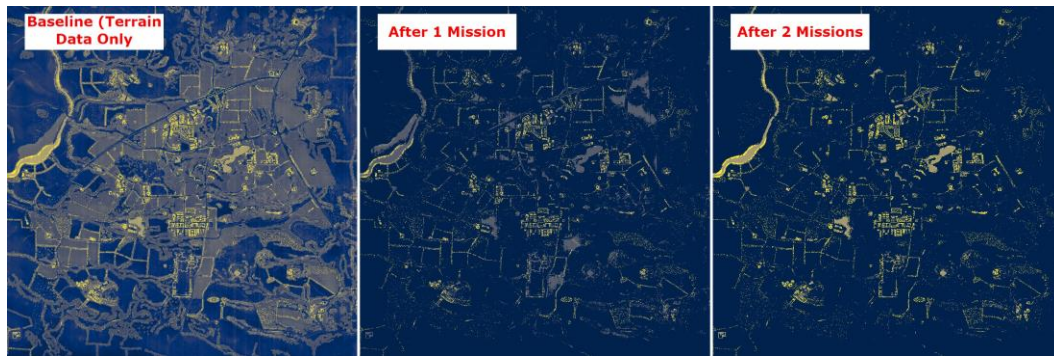


Fig. 11 Comparison of untrained, terrain-only cost map (left), post-Scenario 1 cost map (center), and post-Scenario 2 cost map (right). Differences between these cost maps indicate that the algorithm is indeed learning from each mission and changing its predictions based on participant performance and environmental factors. The extent to which this is happening is subject to study in Phase 2 and beyond.

Using the entire dataset of participant runs, including all their data relating to target identification, mobility obstacles, and AAR ratings, we also trained a single cost

map to display the areas in the entire environment most commonly associated with higher or lower workload predictions. The aggregated predictions are shown in Fig. 12. Areas with obstacles and difficult terrain had higher workload predictions while open, flatter terrain corresponded to lower workload predictions. This helps to validate the learned relationships from the model, especially when comparing these results with the Baseline model in Fig. 7, as these types of terrain would be expected to make the experimental tasks more difficult.

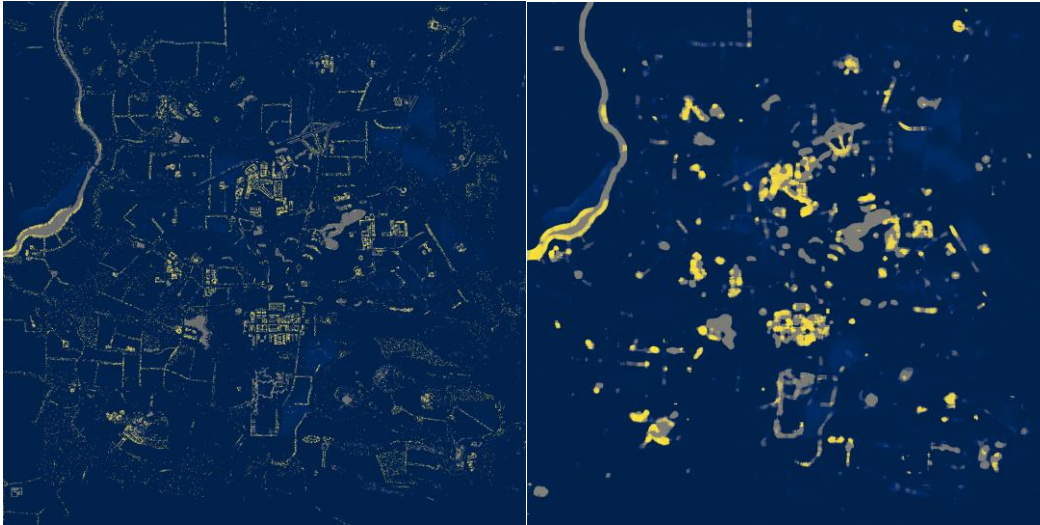


Fig. 12 Final cost maps for the simulated environment, resulting from training the model on the entire participant dataset. On the left, the cost map is represented for each individual pixel. On the right, a percentile filter is applied to “cluster” areas of similar workload, purely for visual representation of the data. Areas involving higher workload predictions appeared to relate most closely to obstacles and difficult terrain, such as the long river system on the left of the map, and the large urban area in the upper center of the map.

2.4 Discussion

The outcomes from this study resulted in several takeaways: a) the algorithm was able to learn and adapt its workload predictions from mission to mission; b) according to participant feedback, the participants used the cost maps to prepare themselves for the upcoming missions; and c) the simulated scenarios were generally appropriate to allow study of the intelligent mission-planning technology. However, the qualitative feedback received from participants identified several areas for further development in the cost map and the scenarios. Following the Phase 1 study, we adjusted the algorithm to make more accurate workload predictions, improved the scenario to include more realistic target-identification tasks (i.e., changing from person-targets to vehicle-targets, as this is primarily a mounted task), gave participants more access to the cost map during the missions to support their planning and decision-making, and expanded the scenarios to involve two participants working together as teammates. We also expanded the

mission planning to consider different types of workloads. Because Phase 1 was a proof-of-concept, we did not allow participants much actual control over how their mission went other than showing them the cost maps and asking them to use that information to prepare themselves for more difficult areas. Thus, another area for improvement in subsequent work involved expanding the planning process.

The primary contribution of the Phase 1 study was in proving the concept of the learning pipeline that ingests information from planning, execution, and AAR steps in the mission process and can adapt predictions from mission to mission. Phase 1 demonstrated that a regression-based algorithm could be used to predict generalized workload across changing mission areas, and that these predictions can be continuously adapted across multiple missions. The intelligent mission-planning capability demonstrated in Phase 1 showed promise for providing mission planners with key information that evolves with dynamic situations and provides predictions that can accelerate the mission-planning process. It is an important step toward our vision of creating systems that efficiently integrate multiple streams of data to coordinate multidomain activities and effects in the dynamic battlespaces of the future. However, given that Phase 1 did not include an experimental design and feedback recommended several areas for improvement, Phase 2 focused on taking further steps toward that vision.

3. Phase 2 Study

3.1 Overview and Goals

The Phase 1 study served as an initial evaluation of the simulator system, experiment design, and intelligent mission-planning technology. Findings from that study have indicated key areas for improvement in the simulation and mission-planning system as we transition the technologies from the one-person task in Phase 1 to a two-person task in Phase 2 and then to eventual operations at platoon-scale and beyond. Some of the Phase 1 findings indicated that a) users wanted to access their planning materials during the mission, instead of only during the planning phase; b) participants often had difficulty identifying the human targets and these targets had less operational relevance to real mounted operations; and c) the mission-planning system needed to include more robust route selection that would better consider the planning materials and previous mission data. For Phase 2, cost maps were made available during the mission, targets were swapped to more operationally-relevant vehicles, and the route selection system was expanded to allow more participant interaction.

The goal of the Phase 2 study was to determine whether the intelligent mission-planning technology improves key measures of the team's performance compared with missions where the technology is not present. The closed-loop pipeline worked in Phase 1 and was capable of evolving its workload predictions from one mission to the next in a simulated situation involving one human with multiple autonomous agents. However, the study lacked experimental comparison and did not allow for much mission planning other than showing participants a cost map and asking them to prepare for areas involving higher workload. For the Phase 2 study, our primary objectives comprised three different areas:

- 1) Expand the Phase 1 study by testing the learning pipeline with a two-person simulation, with each person operating multiple simulated intelligent agents.
- 2) Evaluate whether there were significant differences between mission performance outcomes when the intelligent mission-planning technologies were on or off.
- 3) Further develop the intelligent mission-planning technology and further augment users' capabilities, such as adding levers that allow users to adjust route selection during mission planning to match their preferences, letting users view cost maps during missions, and providing the user with alerts during missions that indicate nearby areas with predicted high workload.

3.2 Methods

3.2.1 Tasks and Stimuli

In this study, each participant/team member acted as the operator of two RCVs that drove through an area while identifying potential targets as friendly or hostile (Fig. 13). The simulation ran on Unreal Engine. The environment was a large area with a mix of rural and urban environments suitable for human–autonomy teaming research.

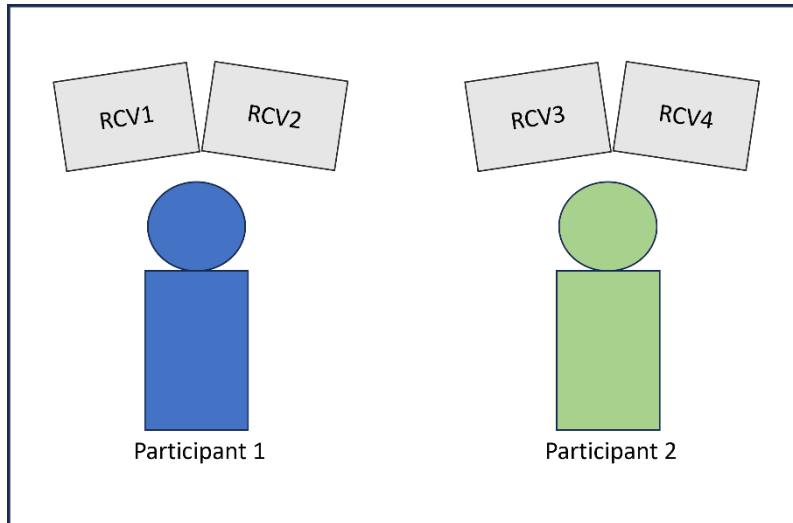


Fig. 13 Participants sat next to each other and each had two computer monitors to use. Each monitor was associated with one RCV.

The format of the study involved one practice scenario, and then four experiment scenarios, each with a planning stage, a mission stage, and an AAR stage. Generally, the study was conceptually similar in design to Phase 1, except with an additional planning capability added to the planning stage. During the planning stage, participants were able to view a refined map of the mission area (with the ability to toggle between cost map and standard views) and the planned start and end points for each of their RCVs (Fig. 14).



Fig. 14 In the planning phase, participants were newly able to cycle between the map view and the cost map view. The cost map color scale was also changed to accentuate differences between low-predicted workload areas (blues, with values close to 0.0), medium-predicted workload (grays with values around 0.5), and high-predicted workload areas (reds with values close to 1.0). This was instituted after reviewing feedback for the Phase 1 study.

Additionally, participants were given a set of sliders they could use, allowing them to select how much weight they wanted the route-generation algorithm to place on the target and obstacle data that feeds into the algorithm's path-planning predictions (Fig. 15). This allowed participants to have some input over the routes their vehicles took during the scenarios. Start and end points were static and predefined during our development of the scenarios, but the paths between those were optimized by the path-planning algorithm based on the user's slider inputs. Our path-planning algorithm used A* to determine the shortest path with the lowest overall costs throughout the path. In our implementation, we set obstacles and extremely steep terrain to have the maximum possible costs, ensuring that A* would not attempt to path through obstacles (such as buildings or fences) or off cliffs. Besides those areas, the rest of the environment's costs generally followed the cost maps generated from each mission by each user.

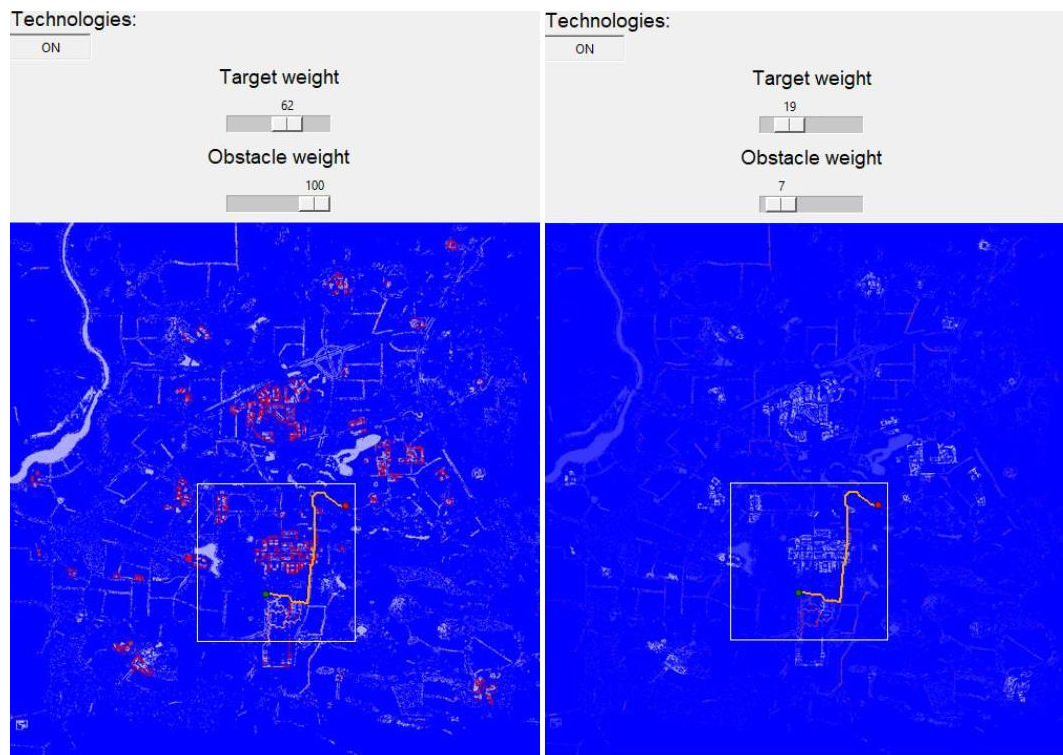


Fig. 15 Slider interface shown during the planning phase. On the left, the user has selected relatively high weights for the target and obstacle cost calculations, resulting in more red areas (high workload) in the cost map predictions. On the right, the user has selected lower weights for the target and obstacle calculations, leading to lower workload predictions (represented in blue). Not shown is the “generate route” button, which the user could press after adjusting the sliders. This would invoke the path-planning algorithm, optimizing the route for that vehicle based on the user's preferred weights. Routes would change based on the weights selected, as the path-planning algorithm optimized for a low overall route cost.

Like in Phase 1, the participants were tasked with monitoring the vehicles while they autonomously navigated through the environment, encountering obstacles and

targets. However, there were several differences implemented for the Phase 2 study that we highlight below. If a certain aspect of the tasks is not mentioned below, the reader can assume it remained the same as in Phase 1.

For the driving task, the four RCVs traveled along paths generated by users during the planning process. This differs from Phase 1, which involved preset paths for all participants. Because paths were not preset, we designed the path-planning algorithm to also inject two to four mobility obstacles along each vehicles' assigned paths, ensuring that all participants would encounter mobility obstacles. However, this was not done with any specific ground-truth logic (hilly terrain has a higher chance to incur an obstacle, flat ground has a lower chance, etc.), and we found that this had the implication of affecting the algorithm's ability to detect and avoid obstacles (this is discussed in further detail in Section 3.4).

For the target recognition task, Phase 2 tasked participants with identifying vehicles, rather than models of humans, as friendly or hostile (Fig. 16).



Fig. 16 Example of a friendly vehicle, the Bradley Fighting Vehicle

The autonomous navigation system operated the vehicles through each of the four scenarios. Cost maps, which helped participants to visually understand where certain parts of the mission may incur a high or low workload, were used to aid the participants during the planning and mission stages during the tech on scenarios (Fig. 17 depicts the mission UI for Phase 2). An automated cueing system was implemented into the interface that sent a sound and text alert to the operator's interface whenever an RCV approached within a certain distance of a high-workload area. This system was only active when the intelligent technologies were on during a Scenario (see Section 3.2.5.1).

The AAR procedure was the same as in the Phase 1 study.

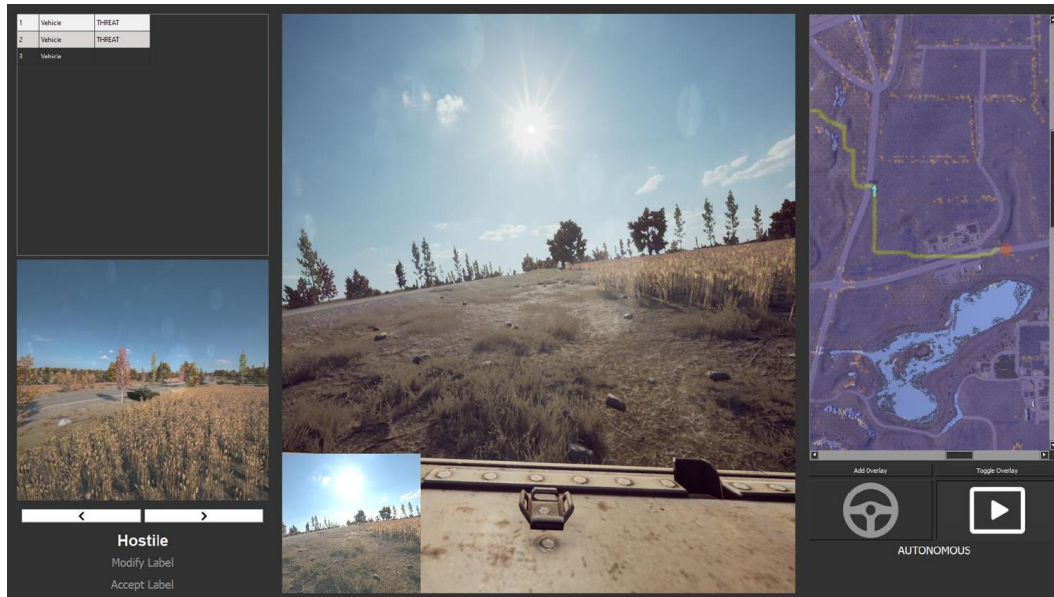


Fig. 17 Sample of the Phase 2 UI. On the left is the target list and a small window that displays thumbnails for detected targets. The center displays the RCV's forward camera. The right side contains a map of the area, and the user can toggle between a regular top-down map overlay or a cost map overlay. The bottom right contains the buttons for switching between manual and autonomous navigation.

3.2.2 Equipment and Setup

The setup for Phase 2 was functionally the same as what was used for Phase 1 except that a second complete workstation setup was added to accommodate the second participant in the study. Both participants still had their own computer, two monitors, steering wheel, etc. See Fig. 5 for a photo of one of the two Phase 2 workstations.

3.2.3 Questionnaires

Participants completed two questionnaires in this study. As with Phase 1, the first was a workload questionnaire given after each mission stage to determine participants' level of workload difficulty they experienced while identifying each target or dealing with each mobility obstacle (Appendix B). The second was the NASA-TLX (Appendix B) that assesses workload on six seven-point scales. Increments of high, medium, and low estimates for each point result in 21 gradations on the scales. This questionnaire was expected to provide a perspective of participants' workload throughout the entire scenario. Finally, at the end of the study, participants completed a study review and feedback questionnaire (Appendix B) aimed at gathering their feedback regarding the design of the scenarios, the simulation, and the intelligent mission-planning technology.

3.2.4 Participants

Fifty-two participants were recruited from the general population around the University of Central Florida area and completed the study. We primarily targeted the student body available on the Sona Systems participant recruitment and management system. The eligibility criteria for participants were as follows. Participants must

- be aged 18–45 years,
- have normal or corrected-to-normal vision,
- have no color vision difficulties,
- be able to view and interact with first-person video games (e.g., driving, shooting) without experiencing motion sickness,
- have no photosensitive epilepsy, and
- have a driver's license. We noted some situations where participants were not able to operate the RCVs using the wheel and pedals, so this condition was added to Phase 2.

As with the Phase 1 study, participants were compensated \$15 per hour of participation, rounded up to the nearest hour, up to a maximum of \$60 for 3.5 h. Participants were paid via Amazon gift cards in \$15 increments.

Because this study required two participants showing up to the lab to be able to conduct the experiment, we instituted an additional procedure that allowed us to recruit three participants per study time slot. If all three participants showed up, we randomly selected one participant to go into what we called “Runoff.”

In Runoff, we allowed the participant to reschedule for our study for another date. If a Runoff participant showed up a second time, they were guaranteed to participate that day (i.e., they would not be put into Runoff a second time). If, due to other participants' no-shows, a Runoff participant was not able to participate that day, they were compensated \$15, thanked, and dismissed. In this way, they were guaranteed to be compensated for their time and efforts in attempting to participate regardless of whether others showed up or not.

If only two out of the three participants showed up, we ran the study as originally planned. If only one participant showed up, we put them into Runoff as described above.

The Runoff procedure was a significant benefit and aided the execution of the study in a timely manner. We experienced a no-show rate of approximately 25%, which

meant that we encountered multiple situations where we would not have been able to run study sessions without this Runoff procedure. Further, in situations where only one participant showed up, the Runoff procedure allowed us to encourage those participants to show up at another date, which helped us fill time slots and run more study sessions across fewer days.

3.2.5 Study Procedure

3.2.5.1 Format and Experimental Conditions

The format of the study involved one training scenario and then four experimental scenarios, each with a planning stage, a mission stage, and an AAR stage. Importantly, the primary manipulation in the study was that for each scenario, the cost map and alerts were either be on (i.e., available to participants during planning and mission), or off (i.e., not shown to participants during planning and mission). Across the four experimental scenarios, teams either had the technologies on-off-on-off or the reverse to counterbalance order and practice effects. These schedules are visualized in Fig. 18. Half of the teams were pseudo-randomly assigned to Group A and the other half to Group B. Thus, the study is a repeated-measures design with a between-subjects component. All teams performed all scenarios in order, the only difference being which scenarios involved the technologies on or off.

Group A					
Scenario	1 (practice)	2	3	4	5
Cost Map	ON	ON	off	ON	off
Group B					
Scenario	1 (practice)	2	3	4	5
Cost Map	ON	off	ON	off	ON

Fig. 18 Explanation of Phase 2 study conditions

The remainder of Section 3.2.5 describes the procedure for each scenario's planning, mission, and AAR stages. Note that if the team was in a scenario where the cost map was on, the teams were able to consult the cost map during the planning stage while deciding which routes for their vehicles to take and were able to view the cost map during the mission stage within the vehicle interface, so they may have been more aware of areas in which to pay greater attention. If the team was in a scenario with the cost map off, those were not available during planning or mission but all else remained the same. When a team had the cost map on for their first mission, a baseline cost map was provided that was generated using

precalculated terrain information (Fig. 4), but the baseline cost map did not consider mission data because no missions were completed by that point.

3.2.5.2 Main Study Procedure

Participants arrived on site, read the informed consent document, and signed it if they intended to participate. They were then shown a set of slides to familiarize them with the missions, their roles, and the controls for operating the vehicles and interacting with the simulation environment.

After finishing the slides, participants completed a simplified practice scenario. They worked together to plan their routes for the training scenario using the route planning interface. Next, they completed the training mission together, in which they identified a few vehicles as friendly or hostile. During the training mission, they became familiar with the layout of the screen and how to monitor and modify the AiTR's classification of targets as friendly or hostile. They also gained experience in operating the RCVs with manual controls and how to change between autonomous and manual modes.

3.2.5.3 Scenario 1 (Practice) Planning Stage

The Scenario 1 planning stage began with teams being shown the map of the environment. Teams were briefed on their goal of monitoring and supporting autonomy in identifying targets and navigating the RCVs. Newly added to the Phase 2 study were the map of the mission area and slider interface, allowing teams to generate route options for each vehicle. Participants chose routes for each of their two vehicles and were able to discuss how they preferred to tackle the mission. Once both participants decided routes for both of their vehicles (covering the four vehicles that the team is in control of), the mission stage begun.

3.2.5.4 Scenario 1 Mission Stage

The missions largely followed the same format as in Phase 1. Each participant encountered approximately two to four obstacles with each of their RCVs, meaning each participant experienced approximately four to eight obstacles per mission. As mentioned previously, these were injected during the path-generation process of the planning stage.

3.2.5.5 Scenario 1 AAR Stage

The Scenario 1 AAR stage began after the team completed the mission. The AAR stepped through the preceding mission and allowed the participant to review each target and indicate how difficult it was to correctly identify it (Appendix B). They did the same for each time they had to take over manual driving. In each of these

instances (target identification and manual driving), participants reported their subjective workload at that time. Images were provided for targets and obstacles to aid their recall. Their responses to these questions were captured by the system to feed the cost map algorithm. After those AAR questions, the participants each completed a NASA-TLX questionnaire.

After the Scenario 1 AAR stage, the data from the previous mission and the AAR were processed by the ML algorithm to update the cost map in case it was available for the next scenario. The experimenter then set up the next scenario as well. During this time, teams were offered a short break.

3.2.5.6 Experimental Scenarios

Scenarios 2 through 5 each occurred in different parts of the map; but the planning, mission, and AAR stages all took the same general format as those in Scenario 1. If the technologies were on for the mission, the participants were able to consult the cost map during planning and the mission and receive alerts throughout the mission. After Scenarios 2, 3, and 4, the team trained the cost map algorithm on data from previous missions and gave participants a short break during that time.

Following Scenario 5, a final questionnaire was presented to collect additional data that may help algorithm development, including asking the participants to explain why they chose the routes that they used. They typed their answers into a text-input document and their responses were saved for later review.

Most study slots lasted approximately 2 h 50 min with an SD of around 25 min. In situations where the study duration reached 3.5 h, the study session was ended by the researcher and the participants were compensated, thanked, and dismissed. Any data collected up to that point was retained. Table 4 provides the procedure and timeline for the study.

Table 4 Procedure and approximate timeline for Phase 2 study

Event	Event time (min)	Estimated cumulative time (h and min)
Informed consent	10	0 h 10 min
Training slides	5	0 h 15 min
Training mission (Scenario 1) and AAR	10	0 h 25 min
Scenario 2 planning	5	0 h 30 min
Scenario 2 mission	10	0 h 40 min
Scenario 2 AAR	10	0 h 50 min
<i>Break, Scenario 3 setup</i>	5	0 h 55 min
Scenario 3 planning	5	1 h 0 min
Scenario 3 mission	10	1 h 10 min
Scenario 3 AAR	10	1 h 20 min
<i>Break, Scenario 4 setup</i>	5	1 h 25 min
Scenario 4 planning	5	1 h 30 min
Scenario 4 mission	10	1 h 40 min
Scenario 4 AAR	10	1 h 50 min
<i>Break, Scenario 5 setup</i>	5	1 h 55 min
Scenario 5 planning	5	2 h 0 min
Scenario 5 mission	10	2 h 10 min
Scenario 5 AAR	10	2 h 20 min
Final questionnaires	5	2 h 25 min
<i>Debrief</i>	5	2 h 30 min

3.3 Results

We computed descriptive statistics for all scenarios and across the entire dataset and also conducted analyses of variance to evaluate whether differences between scenarios or conditions existed. We also calculated route costs by assessing the expected cost associated with each pixel on each RCV’s selected route to observe whether any trends arose. Additionally, qualitative analyses were performed on the feedback questionnaire and observations were recorded from the researchers conducting the study to further future development of intelligent mission-planning technologies. Finally, we also evaluated the cost maps produced by the algorithm, generating per-Scenario aggregate cost maps to assess whether they changed mission to mission.

3.3.1 Quantitative and Descriptives

For Phase 2, we quantified mission performance on four factors: two stemming from target identification and two associated with mobility. Descriptive statistics are reported in Table 5 for each of those factors by condition, and in Table 6 for each of those factors by scenario. There was not a significant difference between means for the key metrics on the on or off conditions. We explore reasons for this in Section 3.4.

Table 5 Descriptive statistics for on and off conditions

Metric	Tech status	Mean (SD)
Mobility RT (s)	On	11.94 (6.04)
	Off	11.28 (4.90)
Mobility travel duration (s)	On	32.41 (9.58)
	Off	33.91 (10.99)
Target accuracy (%)	On	0.86 (0.09)
	Off	0.84 (0.10)
Target RT (s)	On	18.07 (8.96)
	Off	16.23 (6.13)

Table 6 Descriptive statistics by scenario. Table values indicate mean (SD).

Scenario	Mobility RT (s)	Mobility travel time (s)	Target accuracy (%)	Target RT (s)
1 (practice)	17.67 (9.67)	30.85 (11.62)	88.68 (13.26)	36.86 (27.82)
2	12.20 (6.14)	31.47 (10.49)	87.11 (6.80)	21.00 (10.19)
3	12.49 (5.69)	37.40 (18.13)	81.53 (10.32)	18.80 (8.83)
4	11.28 (5.38)	35.29 (9.65)	86.37 (10.82)	15.32 (6.35)
5	10.57 (5.04)	34.52 (10.71)	84.40 (11.40)	14.65 (6.28)

Given the lack of differences in outcome factors between tech on and off conditions, we next moved to assess the effects of those factors across scenarios. Such analysis would at least indicate whether participants improved in performance on those factors after completing several missions.

A similar analysis procedure was carried out to that used for Phase 1. A few outliers beyond 3 SDs were removed from the dataset. Assessments of normality indicated the Phase 2 data did not fit normal distributions and Levene's tests were again mostly significant. As with the Phase 1 data, the Phase 2 data appeared largely nonparametric with unequal variances across factors; thus, we again opted for Kruskal–Wallis H tests to assess whether there were significant differences in the medians of each factor across the five scenarios.

For mobility RT, while the Kruskal–Wallis test indicated a significant difference across scenarios ($H = 9.54$, $p = 0.049$), posthoc pairwise comparisons did not reveal specific scenario pairs with significant differences after Bonferroni correction. This suggests that while there was an overall difference in RTs to mobility obstacles across scenarios, it may not be attributable to any single scenario comparison or that the effect sizes were too small to matter. It appears there was a small, gradual learning effect, but not one significantly notable. Figure 19 represents this, and contains a visible, if small, gradual decrease in mobility RT.

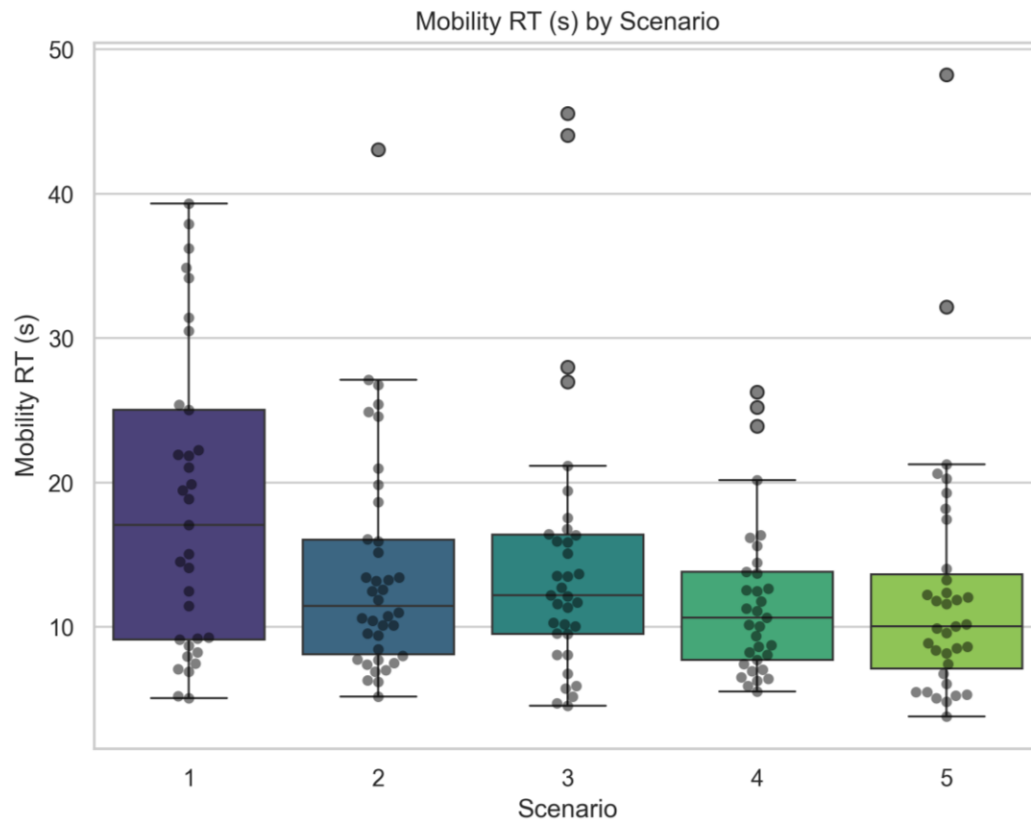


Fig. 19 Mobility RT box plot by scenario, displaying medians and IQRs for the factors. Overlaid swarm plots indicate the overall data distributions for each factor.

For mobility travel time (Fig. 20), no significant differences were observed ($H = 2.93$, $p = 0.57$). Likewise, no significant differences were found between scenarios for the factor target accuracy ($H = 8.26$, $p = 0.08$, Fig. 21). Participants did not get better over time on either factor.

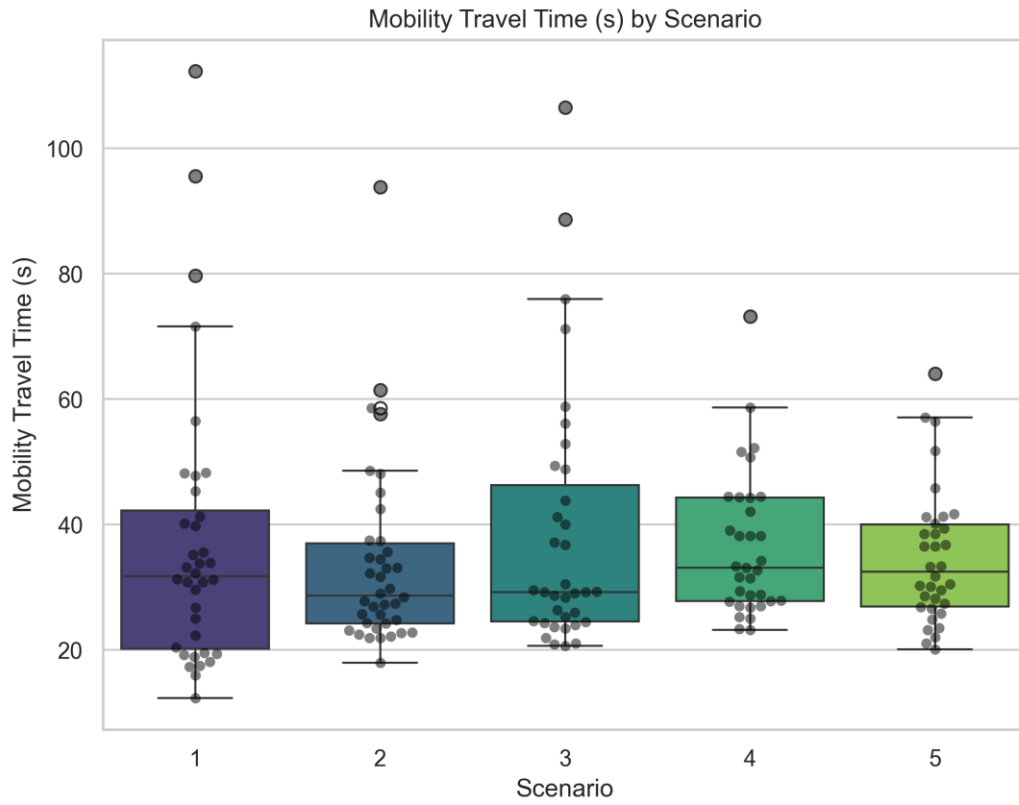


Fig. 20 Mobility travel time by scenario. Scenarios are generally similar-looking with largely overlapping plots, suggesting no improvement was seen in this factor across scenarios.

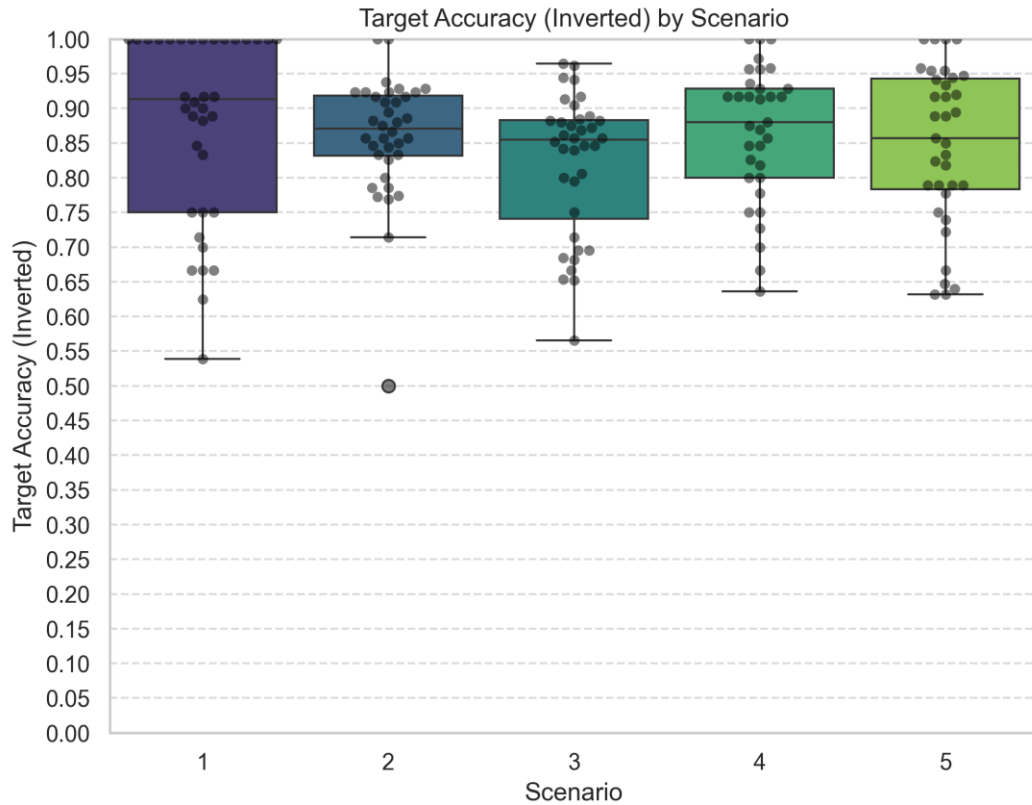


Fig. 21 Target accuracy by scenario. Scenarios are generally similar-looking with largely overlapping plots, suggesting no improvement was seen in this factor across scenarios.

For target RT, a significant difference was observed ($H = 27.94$, $p < 0.0001$). Posthoc Mann–Whitney tests identified differences between several scenario pairs, particularly involving Scenarios 1 and 5, indicating that these scenarios differed markedly in target RTs. This suggests that participants got faster at responding to targets as they completed more missions (Fig. 22).

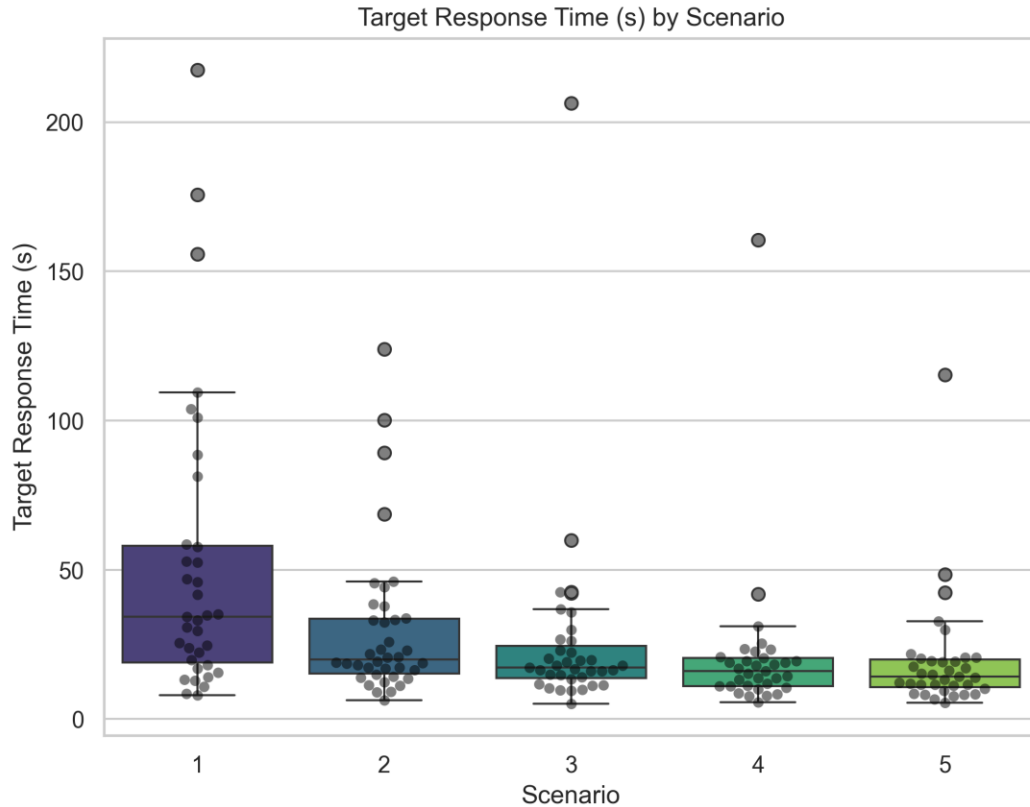


Fig. 22 Target RT by scenario. A downward trend is visually evident and supported by statistical tests, suggesting participants got faster at reacting to targets over the course of several missions. Occasional outliers result from participants sometimes forgetting to respond to a target for several minutes, then coming back to it later. The outliers in this plot were left in as they were found to be genuine mistakes rather than data errors, though some more distant outliers were removed from the plot and analysis.

3.3.2 Qualitative Data

The feedback survey at the end of the study consisted of five questions. Below is a summary of the findings from the questionnaire.

1. What affected the routes you chose during PLANNING phases? In other words, what things did you consider when you chose the routes?

- Most participants chose routes based on ease and distance.
- Factors considered included target weight, number of turns, and obstacle weight.
- Some participants struggled with understanding the planning process and the effect that the sliders had on the route-generation process.

2. How useful was the cost map during the PLANNING phases? Do you feel it aided planning?

- Most found it useful, especially initially.
- Some found it not useful due to not understanding how the sliders affected route plans.
- A few found it somewhat useful, and fewer felt it was not clear what the cost map was trying to convey.

3. How useful was the cost map during the MISSION phases? Do you feel it aided your missions?

- About an even split found it useful and not useful for mission performance.
- Some found the workload alerts helpful, but understanding the alerts was an issue for others.

4. What made the cost map easy or challenging to understand?

- The color scheme was generally helpful in understanding the map.
- Some struggled with the cost map's relevance to routes and understanding the relationship between the slider adjustment and the route plans.
- Suggestions included better explanations and instructional videos.

5. How well do you feel you worked with your teammate? What made your interactions easier or more difficult?

- Most felt teamwork was not essential for the task.
- Some felt they worked well together with good communication and mutual support.
- Observer note: the researchers observed that the majority of participants said little to nothing during scenarios and communicated minimally during the planning phase.

3.3.3 Assessing the Learning Algorithm and Cost Predictions

We next reviewed the ability of the Phase 2 cost map algorithm to evolve mission to mission. We used the entire participant dataset for every scenario along with the baseline cost map that only consisted of predictions stemming from terrain data, resulting in a total of six cost maps that could be reviewed for cross-mission changes.

Figure 23 shows a comparison of the baseline cost map, the post-Scenario 1 cost map, the post-Scenario 2 cost map, and so on, all the way through Scenario 5. As with the Phase 1 study findings, we again noted a visible difference between the

workload predictions across missions, suggesting that some factors relating to the changing missions and participants' performance in them are indeed spurring the algorithm to update its predictions as missions are completed. This suggests that the algorithm and cross-mission learning pipeline are continuing to achieve our goals of providing updated predictions across situations and environments, and the pipeline is performing this cross-mission evolution even in the context of a multi-human multi-agent team.

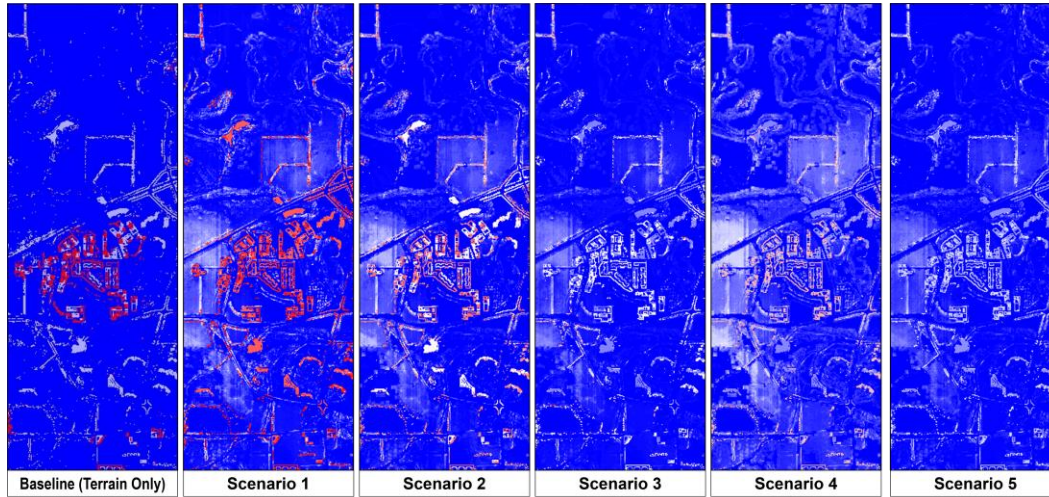


Fig. 23 Comparison of untrained, terrain-only cost map (left) with cost maps generated using data from every subsequent scenario. Differences between these cost maps indicate that the algorithm is indeed learning from each mission and changing its predictions based on participant performance and environmental factors, even as Phase 2 moved the learning pipeline into the multi-human multi-agent context.

Finally, we also assessed the overall costs associated with the path-planning algorithm to verify that the path planner was indeed optimizing for lowering overall route costs. To do this, we conducted a few steps. First, we collected the coordinates (X, Y) for each map pixel along every path generated for every RCV across every scenario. Then, we cross-referenced each pair of pixel coordinates (X, Y) with the cost that the cost map algorithm predicted for that pixel coordinate pair. We repeated that cross-reference for every path pixel traversed by each RCV and then aggregated the costs, resulting in two calculations: 1) the total cost of all routes traversed in each scenario, and 2) the average cost of the routes traversed in each scenario. Those data are represented in Table 7.

Table 7 Total and average path costs for all scenarios

Scenario	Total cost for all paths traversed	Average cost for all paths traversed	Number of RCVs in dataset
2	1507.806	0.048	76
3	1364.189	0.043	72
4	1279.676	0.031	72
5	636.3051	0.017	70

Recall that the goal of the path-planning algorithm is to minimize the overall cost of the route when it looks at the cost map the user generates using their preferred slider weights in the planning process. Thus, we hoped that the average cost for paths would be optimized to be as low as possible and we found this to be the case, as all the average path costs were 0.05 or lower for each scenario. The cost for any pixel coordinate could range from 0 (dark blue on the cost map) to 0.5 (white/gray on the cost map) to 1.0 (bright red on the cost map); thus, the results suggests that the paths were indeed being optimized overall to reduce expected route costs.

3.4 Discussion

The outcomes from Phase 2 have revealed two key findings: 1) the learning algorithm has successfully made adaptive mission-to-mission predictions reflecting participant feedback and behaviors a multi-human, multi-agent environment; and 2) the updates to the intelligent-planning technology via the improved planning process and the UI have resulted in a more effective proof of concept for a multi-human, multi-agent learning pipeline. While our experimental comparison between tech on and off conditions did not reveal significant differences, these key findings demonstrate a promising foundation for continued development and refinement of the learning pipeline as we seek to implement it into larger and more robust simulation environments.

During the review and analysis of the data, we identified several aspects of the experimental design and the learning pipeline that contributed to a reduction in the pipeline’s ability to learn and the experiment’s power. The arrangement of the planning process was as follows: first, the user selected preferred weights for the target and mobility features, which affected how much priority the cost map algorithm placed on those when making its cost predictions. Second, the cost map algorithm calculated costs for every pixel. Third, the path-planning algorithm optimized for a reduced-cost path. Once the user selected that path, the path-generation code for our simulator injected multiple mobility obstacles along the route. This decision was aimed at ensuring that participants would see at least two obstacles per RCV, increasing the likelihood that we would receive more performance data on the mobility obstacles.

However, because the path generator injects the obstacles according to experimenter-set rules and not based on the environmental inputs that the cost map algorithm is trained on, there is not a clear logic about which regions of the environment correlate to increased failure rates for the path-planning algorithm to subsequently get better at avoiding. In other words, the consequence of injecting obstacles into approved routes is that the algorithms are doing everything they can to avoid features related to obstacles, but obstacles are always occurring with a uniform frequency and arbitrary distribution. Thus, there is not a difference in the environment itself that they can learn to avoid. In follow-up work, we are redesigning the environment and tasks so there are innate “ground truth” characteristics of the terrain that contribute to the occurrence of mobility obstacles. For example, urban terrain and areas with dense foliage may have a higher-than-normal likelihood of causing a vehicle to experience a mobility obstacle, while open regions have a lower likelihood. We understand that there is a similar effect for the target-identification process, given that target locations are currently arbitrarily scattered in the environment. Both characteristics will be adjusted in future scenario development to ensure that there is ground truth for the algorithms to identify.

With such ground truth in place, the learning pipeline can then be exposed to routes that may encounter those characteristics, giving it a chance to develop relationships between performance and metrics of workload that reveal problematic terrain, and allow for subsequent mission planning to direct vehicles away from those areas. This adjustment would be more likely to reveal performance differences between tech on and off conditions, and we are adjusting the simulation and algorithms to be able to achieve this.

4. Conclusions and Next Steps

In the future, it will be crucial to support high performance of human–autonomy teams by affording them optimized, intelligent mission-planning capabilities that can synthesize disparate data streams to adapt to complex battlespaces. The Phase 1 study was the first step toward this goal, demonstrating the proof of concept for an algorithm-based workload prediction pipeline that could be used to generalize predictions to an entire map area based on data about the terrain, mission, and AAR. The Phase 2 study improved the design of the pipeline and added additional planning tools for users, expanding the simulated scenario to a multi-human, multi-agent context.

The injection of mobility obstacles according to preset rules after route paths are optimized, rather than by instilling mobility obstacle characteristics into the ground truth of the terrain, limited the algorithms’ ability to learn the world and avoid

problem areas. The overall design of the Phase 2 study and learning pipeline are a promising foundation, but key targets for improvement were identified through the study and will be addressed in our upcoming work to better assess the algorithms' ability to optimize mission plans and provide effective recommendations to users during the planning process. Adjusting those features will allow for subsequent research to more fully explore the extent to which the mission-planning technologies aid performance.

Overall, the Phase 1 and 2 studies have demonstrated that our learning pipeline is capable of modeling workload and performance across missions for multi-human, multi-agent teams, and follow-up work will build on the findings that we have proven thus far. Upcoming work will see the intelligent-planning technologies tested in simulations with more than two humans involved and in more complex, realistic scenarios. The learning pipeline is currently constructed to primarily look for data streams relating to target identification and mobility obstacle handling; however, these are easily adjusted toward whatever mission-relevant data is present. This adjustability is key for our development, ensuring that as the intelligent mission-planning pipeline expands out from the Phase 1 and 2 simulations, it can easily be adapted toward more complex situations. For example, we can set new terrain parameters that are used for the baseline predictions, ingest new streams of data that contribute to cost calculations, and process data in different ways to relay key information to users at more points during the planning, mission, and AAR stages to most effectively support the evolution of teams and the optimization of performance during and across missions.

5. References

- Blickensderfer E, Cannon-Bowers JA, Salas E. Theoretical bases for team self-correction: fostering shared mental models. *Advances in Interdisciplinary Studies of Work Teams*. 1997;4:249–279.
- Brewer RW, Walker AJ, Pursel ER, Cerame EJ, Baker AL, Schaefer KE. Assessment of manned–unmanned team performance: comprehensive after-action review technology development. Paper presented at the International Conference on Applied Human Factors and Ergonomics; 2019; Washington, DC.
- De Dreu CK, Weingart LR. Task versus relationship conflict, team performance, and team member satisfaction: a meta-analysis. *Journal of Applied Psychology*. 2003;88:741.
- Endsley MR. Toward a theory of situation awareness in dynamic systems. *Human Factors*. 1995;37:32–64.
- Hart SG, Staveland LE. Development of NASA-TLX (task load index): results of empirical and theoretical research. *Advances in Psychology*. 1988;52:139–183.
- Li H, Ni T, Agrawal S, Jia F, Raja S, Gui Y, Sycara K. Individualized mutual adaptation in human–agent teams. *IEEE Transactions on Human-Machine Systems*. 2021;51(6):706–714. DOI:10.1109/THMS.2021.3107675
- Marks MA, Mathieu JE, Zaccaro SJ. A temporally based framework and taxonomy of team processes. *The Academy of Management Review*. 2001;26(3):356–376. DOI:10.2307/259182
- Smith-Jentsch KA, Johnston JH, Payne SC. Measuring team-related expertise in complex environments. In: Cannon-Bowers JA, Salas E, editors. *Making decisions under stress: implications for individual and team training*. 1998; American Psychological Association; p. 61–87.
- Wulfmeier M, Rao D, Wang DZ, Ondruska P, Posner I. Large-scale cost function learning for path planning using deep inverse reinforcement learning. *The International Journal of Robotics Research*. 2017;36(10):1073–1087. DOI:10.1177/0278364917722396

Appendix A. Cooperative Agreement Acknowledgment

The research reported herein was sponsored by the U.S. Army Combat Capabilities Development Command (DEVCOM) Army Research Laboratory (ARL) and was accomplished under cooperative agreement W911NF-21-2-0279. The views and conclusions contained in this document are those of the authors and should not be interpreted as representing the official policies or positions, either expressed or implied, of the ARL or the U.S. government. The U.S. government is authorized to reproduce and distribute reprints for government purposes notwithstanding any copyright notation herein.

Appendix B. Questionnaires*

*This appendix appears in its original form, without editorial change.

B.1 Workload Questionnaire (Phases 1 and 2)

While answering the following question for each target or each mobility failure, consider the factors that affected your workload while doing that task. For example, for a target you might consider the target's distance or elevation, how easy it was to see, how overloaded with other tasks you felt at that time, or anything else that may have been relevant to your ability to correctly identify the target. For a mobility failure, you might consider the surrounding terrain, how quickly you noticed the problem, how easy it was to rectify, or anything else that may have been relevant to your ability to respond to the mobility failure.

(1: Very Low, 2: Low, 3: Slightly Low, 4: Neutral, 5: Slightly High, 6: High, 7: Very High)

What is the level of workload you experienced while identifying this target?

1 2 3 4 5 6 7

What is the level of workload you experienced while dealing with this mobility failure?

1 2 3 4 5 6 7

B.2 NASA-TLX (Phases 1 and 2)

What is your Participant Number? _____

What Scenario did you just complete? _____

Mental Demand

How mentally demanding was the task?



Physical Demand

How physically demanding was the task?



Temporal Demand

How hurried or rushed was the pace of the task?



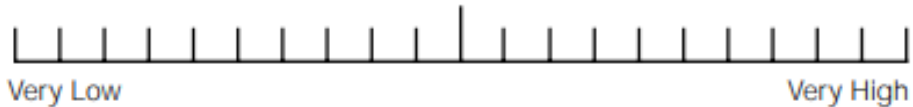
Performance

How successful were you in accomplishing what you were asked to do?



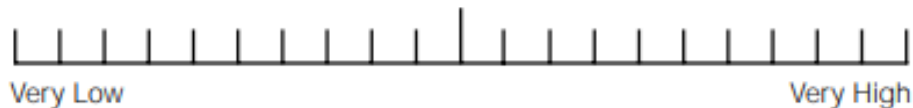
Effort

How hard did you have to work to accomplish your level of performance?



Frustration

How insecure, discouraged, irritated, stressed, and annoyed were you?



B.3 Driving Experience Questionnaire (Phase 1)

Do you have a driver's license? Yes / No

If Yes, which of the following best represents approximately how often you drive?

Daily A few times a week A few times a month A few times a year

Approximately how many years have you been driving? _____

B.4 Study Feedback Questionnaire (Phase 1)

The following free-response questions were administered to participants via computer during the AAR that followed the final mission.

1. Why did you select the route that you chose in Scenario 3?
2. Did you feel like the missions reached an appropriate level of difficulty? If any aspects were too easy or difficult, describe those aspects.
3. How useful was the cost map for understanding the upcoming mission?
4. What would you change about the cost map?
5. What made the cost map easy or challenging to understand?
6. Did you have significant difficulty identifying some targets? If so, what aspects contributed to the difficulty?

B.5 Study Feedback Questionnaire (Phase 2)

The following free-response questions were administered to participants via computer after completing the final Scenario.

1. What affected the routes you chose during PLANNING phases? In other words, what things did you consider when you chose the routes?
2. How useful was the cost map during the PLANNING phases? Do you feel it aided planning?
3. How useful was the cost map during the MISSION phases? Do you feel it aided your missions?
4. What made the cost map easy or challenging to understand?
5. How well do you feel you worked with your teammate? What made your interactions easier or more difficult?

List of Symbols, Abbreviations, and Acronyms

AAR	after-action review
AiTR	aided target recognition
ARL	Army Research Laboratory
DEVCOM	U.S. Army Combat Capabilities Development Command
FY23	fiscal year 2023
FY24	fiscal year 2024
IQR	interquartile range
ML	machine learning
NASA	National Aeronautics and Space Administration
RCV	robotic combat vehicle
RT	response time
SD	standard deviation
TLX	Task Load Index
UI	user interface

1 DEFENSE TECHNICAL
(PDF) INFORMATION CTR
DTIC OCA

1 DEVCOM ARL
(PDF) FCDD RLB CI
TECH LIB