



ARL-TR-10311 • MAR 2026



Building Shared Situational Awareness across Heterogeneous Units for Improved Squad Lethality

by Chloe Callahan-Flintoft, Joyce Tam, Osben Toulson, Richard Diego, Kevin King, Justin Brody, Joseph Conroy, Andrew Tweedell, Thomas Rohaly, Harrison Crowell, Heather Roy, and Jonathan Touryan

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.



Building Shared Situational Awareness across Heterogeneous Units for Improved Squad Lethality

**Chloe Callahan-Flintoft, Joseph Conroy, Andrew Tweedell,
Harrison Crowell, Heather Roy, and Jonathan Touryan**
DEVCOM Army Research Laboratory

Joyce Tam
Oak Ridge Associated Universities

Osben Toulson, Richard Diego, Kevin King, and Thomas Rohaly
DSC Corporation

Justin Brody
Fibertek, Inc.

REPORT DOCUMENTATION PAGE

1. REPORT DATE		2. REPORT TYPE		3. DATES COVERED					
March 2026		Technical Report		<table border="1" style="width: 100%; border-collapse: collapse;"> <tr> <td style="width: 50%;">START DATE</td> <td style="width: 50%;">END DATE</td> </tr> <tr> <td>9/1/2024</td> <td>6/30/2025</td> </tr> </table>		START DATE	END DATE	9/1/2024	6/30/2025
START DATE	END DATE								
9/1/2024	6/30/2025								
4. TITLE AND SUBTITLE									
Building Shared Situational Awareness across Heterogeneous Units for Improved Squad Lethality									
5a. CONTRACT NUMBER		5b. GRANT NUMBER		5c. PROGRAM ELEMENT NUMBER					
5d. PROJECT NUMBER		5e. TASK NUMBER		5f. WORK UNIT NUMBER					
0602184A/CN2									
6. AUTHOR(S)									
Chloe Callahan-Flintoft, Joyce Tam, Osben Toulson, Richard Diego, Kevin King, Justin Brody, Joseph Conroy, Andrew Tweedell, Thomas Rohaly, Harrison Crowell, Heather Roy, and Jonathan Touryan									
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES)				8. PERFORMING ORGANIZATION REPORT NUMBER					
DEVCOM Army Research Laboratory ATTN: FCDD-RLA-FC Aberdeen Proving Ground, MD 21005				ARL-TR-10311					
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: Chloe Callahan-Flintoft, 0000-0002-1874-2032; Joyce Tam, 0000-0003-0020-3177; Andrew Tweedell, 0000-0002-8182-7768; Harrison Crowell, 0000-0001-8763-8733; Heather Roy, 0000-0002-7843-2335; Jonathan Touryan, 0000-0001-7343-7869; Kevin King, 0000-0003-4888-7491; Thomas Rohaly, 0000-0002-1623-3015									
14. ABSTRACT									
Human-machine integration (HMI) requires a shared understanding of context; that is, what are the mission objectives, and what are the environmental parameters that might affect execution of the mission? To address this need, we focused on building shared understanding of the surroundings by compiling gathered visual information during a mission-relevant scenario, hereafter referred to as situational awareness (SA). Humans build SA through a collection of cognitive mechanisms that allow them to attend or ignore information depending on the parameters of the task and the environment. Sharing this information with autonomy could enable heterogenous human-machine units to respond adaptively in dynamic real-world mission scenarios. However, to do so requires human SA to be characterized in a machine-readable format. Toward that aim, in the current work, two Soldiers equipped with HoloLens headsets and sensor-augmented rifles searched a mock urban environment for targets. A theory-grounded, neural model of human visual attention provided a quantitative characterization of the information in each Soldier's visual field most likely to be actively perceived and remembered by the Soldier. To amass this information across multiple human members of the squad as well as other robotic systems, we built the networking capabilities necessary to record synchronized multimodal data streams and identified challenges in localizing these data streams into a shared coordinate system. Finally, we discuss lessons learned, remaining gaps, and suggestions on how future work could close those gaps to achieve bidirectional, adaptive HMI.									
15. SUBJECT TERMS									
Humans in Complex Systems, computational cognitive modeling, situational awareness, human-machine integration, cognitively informed adaptive autonomy									
16. SECURITY CLASSIFICATION OF:				17. LIMITATION OF ABSTRACT					
a. REPORT		b. ABSTRACT		c. THIS PAGE					
UNCLASSIFIED		UNCLASSIFIED		UNCLASSIFIED					
				UU					
18. NUMBER OF PAGES									
35									
19a. NAME OF RESPONSIBLE PERSON				19b. PHONE NUMBER (Include area code)					
Chloe Callahan-Flintoft				(520) 691-5090					

STANDARD FORM 298 (REV. 5/2020)

Prescribed by ANSI Std. Z39.18

Contents

List of Figures	v
1. Introduction	1
2. Background	2
2.1 Modeling Human SA	2
2.2 Localization of the Environment	4
2.3 Adaptive Autonomy Informed by Human Context	4
3. Methods	5
3.1 Vignette	5
3.2 Cognitive Model	6
3.2.1 Modifications to the RAGNAROC Model	6
3.2.2 Model Evaluation	7
3.3 Simulated UAS	8
3.4 Network	9
3.4.1 HoloLens	10
3.4.2 Custom Instrumented Rifles	11
3.5 Inferred Shared Coordinate Space	12
3.6 Real UAS	13
4. Results	14
4.1 Model	14
4.2 Simulated UAS	16
4.3 Inferred Shared Coordinate Space	17
5. Discussion	17
5.1 Outstanding Challenges and Future Solutions	18
5.1.1 Translating Retinotopic 2D AM Activation to a 3D World Map	18
5.1.2 Developing and Testing Alternative Adaptive UAS Flight Algorithms	19

5.1.3	Raising the Signal-to-Noise Ratio in the Inferred Shared Coordinate System	20
5.1.4	Using Opportunistic Sensing to Improve UAS Target Classification	20
5.1.5	Improving Hardware and Software Support for Real-World Multi-Human Team Data Collections	21
6.	Conclusion	22
7.	References	23
	List of Symbols, Abbreviations, and Acronyms	25
	Distribution List	27

List of Figures

- Figure 1. (A) RAGNAROC’s design features a hierarchical neural network that processes visual information. Signal strength within this network is dynamically adjusted by both saliency and task relevance. The interplay of these factors generates a priority map; the AM represents the distribution of attentional allocation across the visual field. (B) Lateral interactions within the AM are mediated by inhibitory-gating (IG) nodes, creating a mechanism for surround and reactive suppression. (i) Active AM neurons stimulate adjacent IG nodes using a difference-of-Gaussian pattern. (ii) Each IG node then inhibits a corresponding AM neuron. (iii) An AM neuron, excited over the high threshold, can, in turn, suppress its corresponding IG node. 3
- Figure 2. Schematic of operation and data collection. Soldiers (demarked as rifle [RFL] 1 and RFL 2) were tasked to move through the MOUT site from point A to point B, shooting red boxes as they found them. Both Soldiers were outfitted with augmented rifles and HoloLens headsets. Red squares depict example locations of targets. A UAS was also used to search the MOUT, flying a preprogrammed, grid search pattern. 5
- Figure 3. Example POV video frame. Green bounding boxes indicate object detection output from Grounding DINO. Overlaid heat map indicates attentional priorities assigned by RAGNAROC. 7
- Figure 4. Network architecture schematic. Notes: EO = electro-optical; IMU = inertial movement unit. 10
- Figure 5. Instrumented rifle with video, IMU, GPS, and magnetometer. 11
- Figure 6. Example images from taken from the UAS at lower (left) and higher (right) altitudes. YOLO classification boundary boxes have been superimposed on images to show where objects (boxes) were detected. 14
- Figure 7. Gaze error produced by RAGNAROC, GBVS, and random selection. Error bars and markers indicate quantiles and mean values of individual Soldiers. 15
- Figure 8. Attentional priority (Z-scores) in period before target engagement via rifle trigger pull. Golden line represents priority at the to-be-engaged target’s location, and tan line represents priority at the location of a control target that was also visible..... 16
- Figure 9. Squad-level coverage of the environment. Left: Visualization of projected field of regard for both Soldiers in orange and green, respectively. Dotted lines denote previous locations of either Soldier. Field of regard is plotted as a 70° wedge of rays, originating from the Soldier’s head location and rotated using the Soldier’s head orientation. Rays terminate at objects (e.g., buildings, defined by an edge detection algorithm run on the map image) and increase in transparency over time. Right: Squad-level coverage of the environment plotted across time to demonstrate the effect of

incorporating UAS support with independent vs. adaptive programming of flight speed.....	16
--	----

Figure 10. Results of inferred shared coordinate space. A. Confidence heat map. Warmer colors denote areas that were frequently predicted to contain targets across video frames. Red squares are ground truth locations of targets (red boxes). Yellow squares are ground truth locations of non-targets (yellow). Note that other non-targets were present in the environment, but locations were not stored. Green squares denote locations of estimated targets by placing a green square directly in front of the Soldier’s head using the head pitch to estimate depth at the time of trigger pulls. Solid circles around target denote the smallest radius used in the radial sweep (20 pixels), and dotted circles represent the largest radius used (70 pixels). B. Precision, Recall, and F1 score.	17
--	----

1. Introduction

Human-machine integration (HMI) is a critical and evolving challenge for U.S. Army modernization efforts as the artificial intelligence (AI) application space grows. To capitalize on the potential capabilities enabled through AI, specific attention must be paid toward the interaction between humans and autonomous systems, to create a seamless integration between the two. Specifically, effective integration requires that the human's role within this human-machine unit does not create a processing bottleneck and that the machine provides intuitive aid or information that integrates into the human's natural processing of the environment. One major component of HMI is being able to sense and infer information about the state of the human and the environment without creating additional burden to the human, a concept often referred to as opportunistic sensing (Lance et al. 2020). This is a critical component in moving human-machine units from interactions that require the human's full attention at the cost of all other tasks to interactions in which humans can allocate their attention towards critical priorities while the autonomous system continually infers human intent and mission context and provides assistance accordingly.

To achieve this objective, the Distributed Information for Enhanced Squad Lethality (DIESL) program seeks to develop novel approaches and technologies to infer intent and mission context through passive observations of human behavior. The work presented here specifically focused on how to characterize a Soldier's situational awareness (SA) and translate that awareness into a machine-readable format by which to inform an adaptable autonomous system. Data was collected from two Soldiers, outfitted with HoloLens headsets as well as rifles augmented with various sensors, as they conducted military scenarios at a military operations on urbanized terrain (MOUT) site used for research. In addition to the sensors on the Soldiers and the rifles, the Soldiers were paired with a Red Cat Teal 2 Drone unmanned aerial system (UAS) to serve as a scout, to build a shared human-machine understanding of the environment and improve search metrics (e.g., coverage of the area and target acquisition). This report details the various efforts that supported this objective, including cognitive modeling, network integration, and algorithm development. The rest of this report is divided into the following sections: 1) background on each component; 2) methods; 3) initial results; and 4) discussion on lessons learned, remaining problems to solve, and proposed future directions.

2. Background

SA can have a variety of meanings in a military context (Endsley and Garland 2000). For the purposes of the current work, we define SA as a general understanding of the surrounding physical environment. For the Soldier, we focus on the visual modality of SA, seeking to characterize what visual information was available to the Soldier (i.e., fell within their field of regard) as well as what information was actively attended (i.e., more likely to be perceived and remembered and could be responded to with enhanced speed and accuracy). For the human–autonomy unit, SA not only includes what the Soldiers see and attend to, but also a shared coordinate space by which to understand the spatial relationships among squad members, targets, and geographic landmarks. Each of the components described in this report (i.e., cognitive modeling, inferred coordinate space, and adaptive HMI) is a demonstration of the potential gains that could be made in moving these respective fields beyond the current state of the art. For example, currently, very few theory-grounded models for predicting human SA are designed for or tested in full ambulatory exploration of an outdoor environment, which are the very conditions necessary to make them useful to Warfighter technologies. Similarly, integrating multimodal data streams from multiple unit members into a shared coordinate space is an outstanding challenge, especially in a military environment, where sensors must go undetected by adversarial agents. Both the cognitive and engineering challenges outlined in this work aim to push both fields forward by using the combined effort to generate shared human–machine SA. Finally, this work demonstrates how such an SA map, integrated across a heterogeneous unit, could be used to enhance efficiency and lethality by enabling adaptively responsive autonomous systems.

2.1 Modeling Human SA

The human brain creates sensory representations of its environment adaptively to maintain a balance between representational fidelity and resource optimization. Human SA is formed from these internal representations whereby information deemed important based on immediate cues or established experience is prioritized with selective attention. The formed representation at any given moment is therefore an adaptive (nonidentical) copy of the available external input. The crux of modeling human SA lies in 1) leveraging scientific understanding of how the human brain performs such transformations, 2) simulating the resulting representations, and 3) incorporating this knowledge in the optimization of HMI.

RAGNAROC (Reflexive Attention Gradient through Neural AttRactOr Competition; Wyble et al. 2020) is a rate-coded neural model that simulates the

deployment of covert, reflexive attention (Figure 1). The model has two adjustable parameters: “bottom-up” saliency, representing visual conspicuity (Li 2002), and “top-down” relevance, representing feature importance. Taken together, these parameters capture key aspects of attentional modulation. Visual information flows through a hierarchy of retinotopically organized neurons to generate the attention map (AM). Neural activity is modulated by both saliency and relevance, shaping competitive dynamics within the network. The resulting AM is a dynamic, retinotopic map that simulates attentional allocation across the visual field, indicating which locations are prioritized and which are suppressed.

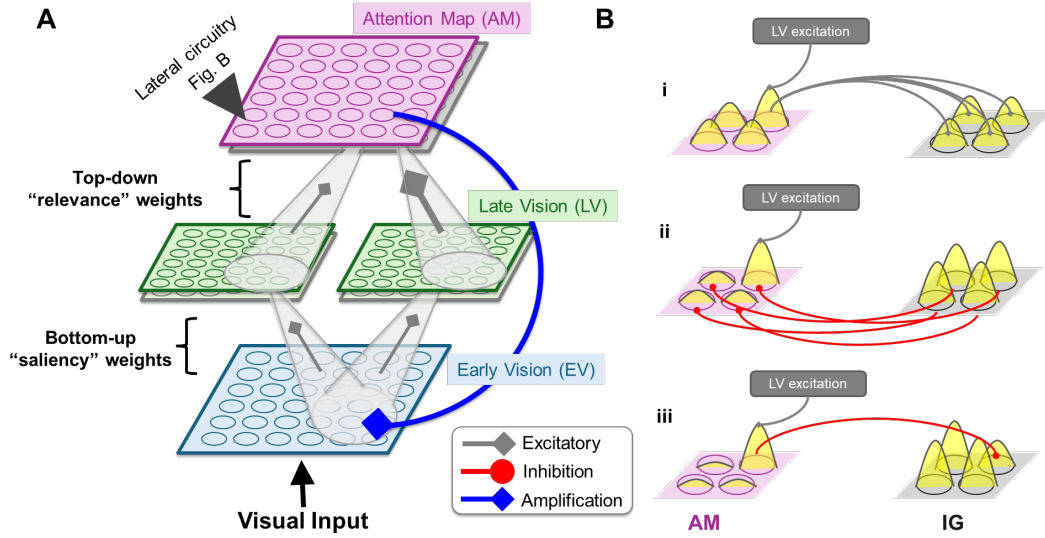


Figure 1. (A) RAGNAROC’s design features a hierarchical neural network that processes visual information. Signal strength within this network is dynamically adjusted by both saliency and task relevance. The interplay of these factors generates a priority map; the AM represents the distribution of attentional allocation across the visual field. (B) Lateral interactions within the AM are mediated by inhibitory-gating (IG) nodes, creating a mechanism for surround and reactive suppression. (i) Active AM neurons stimulate adjacent IG nodes using a difference-of-Gaussian pattern. (ii) Each IG node then inhibits a corresponding AM neuron. (iii) An AM neuron, excited over the high threshold, can, in turn, suppress its corresponding IG node.

RAGNAROC represents a paradigm-general understanding of visual attention by simulating a variety of empirical benchmarks involving visual search time, accuracy, and neural correlates. It also demonstrates capabilities to simulate multiple performance metrics—including gaze pattern, target selection, and sudden-onset hazard detection—in an immersive virtual reality (VR) environment (Tam and Callahan-Flintoft 2025). In this report, we enhanced the model’s compatibility with naturalistic visual input by creating an interface between it and the Graph-Based Visual Saliency (GBVS) model (Harel et al. 2007), which

evaluates an image’s visual conspicuity and the Grounding DINO model (Liu et al. 2024), which performs zero-shot object detection.

2.2 Localization of the Environment

To build shared SA between multiple humans and autonomous systems, a common coordinate space by which to understand the environment is required. This is a non-trivial challenge and an active area of research and technology development (e.g., simultaneous localization and mapping). However, this is especially challenging in a military context because using technology such as LiDAR (i.e., light detection and ranging, which can help build a shared spatial representation of the environment) can leave the Soldier vulnerable to enemy detection (Wang et al. 2024). Developing alternative, non-emissive methods of sensing and inferring a shared representation of the environment then becomes imperative. One potential source for inferring location information from Soldiers is point-of-view (POV) video data. Knowing what a person can see as they move through the environment can give insights into the spatial relationships of objects in the environment, but the translation of 2D video frames into a 3D map is difficult. A component of this process, detecting objects in each frame, is a well-investigated field with many computer vision algorithms available (Szeliski 2022). However, the tools for relating those detected objects to an x,y,z coordinate space are less robust. In this report, we use the POV video in conjunction with the Soldier’s head orientation and location to estimate where visible targets were located on a shared coordinate space.

2.3 Adaptive Autonomy Informed by Human Context

The performance of human–autonomy units can be improved through opportunistic sensing of human context, wherein the autonomous agent senses human and environmental signals and leverages this information to adapt to dynamic operational scenarios. For example, in a case where UASs are used to support humans in search and navigation, various parameters of the drone’s operation, including its flight path, speed, and altitude, can be adaptively set based on human-derived context. In this report, we performed a series of simulations that incorporated the Soldiers’ combined field-of-regard to adapt the search path of a UAS and evaluated the effectiveness of this approach in enhancing squad-level search efficiency.

3. Methods

3.1 Vignette

This data collection was conducted at the U.S. Army Combat Capabilities Development Command (DEVCOM) Army Research Laboratory's (ARL's) Robotics Research Collaboration Campus (R2C2) located in Graces Quarters, Maryland. Two Soldiers were tasked with searching the MOUT facility. Cardboard boxes that were painted a variety of colors were scattered through the MOUT site. The boxes were intentionally clustered in the areas between buildings (Figure 2). In total, there were 14 green, 17 blue, 22 red, 20 yellow, 7 orange, and 8 pink boxes. The Soldiers were told to search for and shoot all red boxes. They were instructed to navigate from one side of the MOUT site to the other using a given route, but otherwise they could use whatever tactical strategies were warranted.

A ground truth dataset was acquired through a separate same-day, data collection with the HoloLens. To collect this dataset, an operator stood over each yellow (non-target) and red (target) box, thereby tagging the location of the boxes within the headset's spatial reference frame. These two box types were chosen to give a rough approximation of the surrounding environment because recording all the boxes' locations was not feasible given the number of boxes used to create visual clutter.

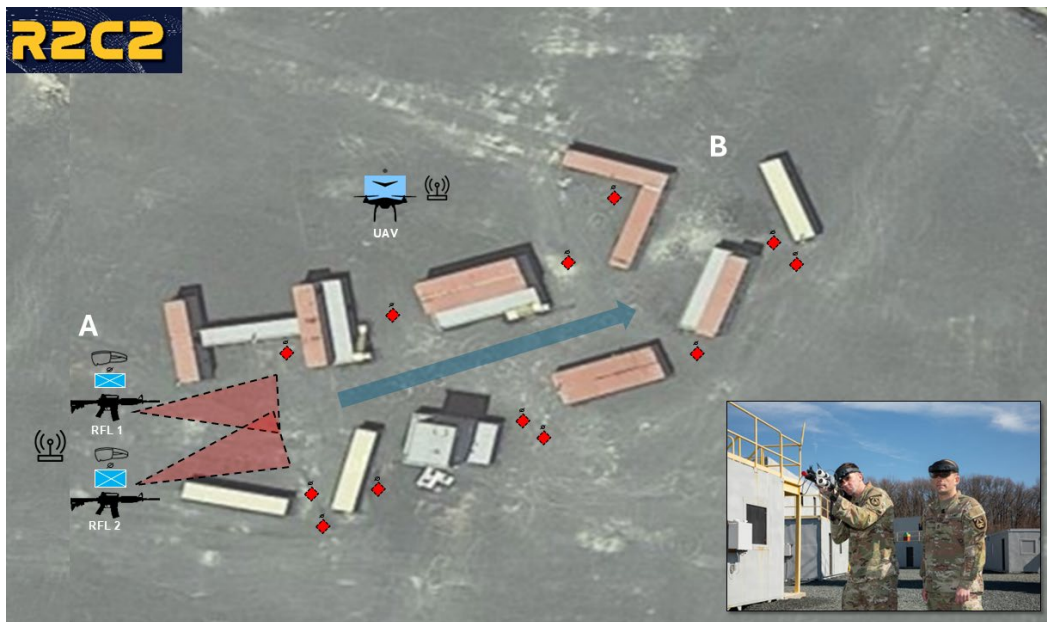


Figure 2. Schematic of operation and data collection. Soldiers (demarked as rifle [RFL] 1 and RFL 2) were tasked to move through the MOUT site from point A to point B, shooting red boxes as they found them. Both Soldiers were outfitted with augmented rifles and HoloLens headsets. Red squares depict example locations of targets. A UAS was also used to search the MOUT, flying a preprogrammed, grid search pattern.

3.2 Cognitive Model

To estimate attentional deployment in a real-world, first-person viewpoint video, several adjustments needed to be made to RAGNAROC. In the original version of the model, attentional deployment to a scene is estimated by providing the retinotopic x,y location of all the objects in that scene and assigning each object both a saliency weight as well as a task-relevancy weight. The choice of input was a necessary simplification to focus development efforts on the attentional mechanisms at the upper layers of the model. However, this sort of input is incompatible with real-world video data, because it fails to account for the complexity of the natural environment. To remedy this, the following modifications were made to the model.

3.2.1 Modifications to the RAGNAROC Model

Assigning a fixed value as an object’s saliency in real-world video data is problematic. For example, many elements of a scene (e.g., the sky and ground) are not typically thought of as objects but rather background components, because they are often large and discontinuous to the viewer. However, these components have texture and offer visual information that should be available to the model. Another reason that fixed, object-based saliency values do not work in this context is that as a person moves through the environment, their perspective and lighting changes, which has considerable impact on the saliency of objects in view. This was not a consideration during the original construction of the model because it was based on 2D, desktop display experiments. Therefore, to accommodate the additional complexity of the dynamic video input, the early visual layers were replaced with the GBVS model (Harel et al. 2007). GBVS is a popular saliency model that identifies visual conspicuity in a variety of feature dimensions, and it was inspired by the neural mechanics of early visual cortex, making it an appropriate stand-in for the early processing layers of RAGNAROC.

Another challenge lies in locating the target objects such that suitable task-relevancy weightings could be assigned to the data. As the Soldier navigates the site, an object’s relative position in the first-person video data will vary according to the Soldier’s movements, despite the object being stationary in the world. Therefore, before feeding the data to RAGNAROC, we performed object detection to recover the objects’ relative positions at each video frame, so that a higher task-relevancy weighting could be assigned to target locations. For object detection, we applied the Grounding DINO model (Liu et al. 2024) to the HoloLens POV video. Grounding DINO performs detection using human input that describes the object category (e.g., “red box”) and attempts to detect and locate the objects with confidence-rated bounding boxes (Figure 3). Furthermore, Grounding DINO is

currently the benchmark of zero-shot object detection, meaning that it can reliably detect object categories outside of its training set. Successful integration of Grounding DINO into the simulation pipeline will allow this work to be scalable to a variety of combat scenarios involving various target categories. Aerial data was similarly post-processed but leveraged few-shot training and a fine-tuned You Only Look Once (YOLO) v12 (Tian et al. 2025) model due to the relatively small size of the boxes as perceived from the UAS.



Figure 3. Example POV video frame. Green bounding boxes indicate object detection output from Grounding DINO. Overlaid heat map indicates attentional priorities assigned by RAGNAROC.

Finally, the temporal continuity of the video data must be considered. In a typical desktop experiment, visual events are presented on a trial-by-trial basis. When RAGNAROC predicted attentional deployment in those experiments, the model’s start and stop times could simply be set according to the trial’s start and stop times. However, these temporal boundaries do not exist in the current real-world video data but instead had to be estimated. Moment-to-moment viewpoint change (defined as rotational velocity from the sum of head and gaze vectors) was used such that an above-threshold viewpoint change would trigger a “reset,” like a trial-onset indicator. The viewpoint threshold in the current work (set to trigger resets when gaze moved over 0.1 degrees of visual angle [dva], result in a reset in 96.21% of the frames) was determined in previous VR work (Tam and Callahan-Flintoft 2025) in which a range of threshold values was tested with a larger dataset, wherein threshold selection was based on the model’s capability in predicting gaze locations.

3.2.2 Model Evaluation

With saliency weightings computed from GBVS and task-relevancy weightings assigned with reference to object detection results from Grounding DINO,

RAGNAROC processes the video data frame-by-frame, outputting an x-by-y-by-t data structure corresponding to the model’s simulation of attentional priority at each x,y location of the visual field and how it evolved over time (t). The validity of RAGNAROC’s predictions was evaluated in two ways. First, we tested whether the model could predict the Soldiers’ gaze pattern by computing the gaze error, i.e., the Euclidian distance between the location assigned maximal priority by RAGNAROC, and the projected gaze location determined by the HoloLens for each video frame. RAGNAROC’s gaze error was compared with GBVS’s, which did not account for target object locations or the neural mechanisms embedded in RAGNAROC, as well as a random location baseline. Gaze error was reported in estimated dva with the assumption that the width of the video frame corresponded to 110 dva.

The second evaluation metric regarded how well RAGNAROC predicted upcoming target engagement via rifle trigger pulls. Each time a Soldier pulled the trigger, the timing was recorded and video frames up to (~3.3 s preceding) and including the time of target engagement were analyzed. Predicted attentional priority at the to-be-engaged target’s location was compared with that at another visible target’s location (control). Note that RAGNAROC was set to upweight target locations using labels from Grounding DINO, but not all target instances were identified through that method. Therefore, for this evaluation, RAGNAROC was rerun with semi-manual target labels replacing Grounding DINO labels in the pre-engagement periods. Specifically, the coordinates of both the to-be-engaged and the control targets were first manually labeled at the time of engagement. Interest points within each bounding box were then identified using the minimum eigenvalue algorithm (Shi 1994) and then fed to the Kanade–Lucas–Tomasi (Lucas and Kanade 1981; Tomasi and Kanade 1991) feature tracker to estimate the object’s locations across the pre-engagement period. A video frame was included only if both targets (to-be-engaged and control) were in view.

3.3 Simulated UAS

The effectiveness of incorporating opportunistic sensing into the control of UAS search algorithm was evaluated by comparing the level of squad-level coverage efficiency achieved by including UAS support that was adaptive (with opportunistic sensing) or independent (without opportunistic sensing). In both conditions, the simulated UAS provided coverage support following a grid search pattern. While the independent UAS progressed on its flight path at constant speed, the adaptive UAS sensed the dynamics of spatial coverage achieved by human Soldiers on the ground and adaptively adjusted its flight speed grid-by-grid, such

that relatively more time was allocated to grids with less existing on-ground coverage.

To estimate on-ground coverage achieved by Soldiers throughout the search, image analysis was first performed on a top-down view of the MOUT site to identify “ground” areas, as opposed to unreachable areas such as walls and ceilings. Each Soldier’s location and heading (recorded with HoloLens) was then used to compute their field of regard at any given moment. Percentage on-ground coverage was given by the ground area included in at least one Soldier’s field of regard divided by the total ground area.

The time an adaptive drone uses to cover a grid, t_A , is equivalent to the baseline time spent by an independent drone, t_I , except that the adaptive drone speeds up as the percentage of on-ground coverage already achieved by human Soldiers, w , increases:

$$t_A = \max(t_{min}, t_I - \ln(w + 1)) , \quad (1)$$

where t_{min} is a positive constant referring to the minimum amount of time required for the drone to fly through a grid.

3.4 Network

The network architecture for this data collection had to support two HoloLens head-mounted displays (HMDs), two rifles, and the UAS data streams. To do this, the robotic operating system (ROS) enclave was used. For the HoloLens, this involved sending a heartbeat to time sync across devices via a laptop. Data from the HoloLens was recorded on each device and then downloaded on the laptop after each run. The rifles sent real-time data wirelessly to the ROS enclave (Figure 4). UAS data was saved separately to the ROS enclave as a UAS search was performed sequentially, rather than simultaneously, with human search due to safety regulations.

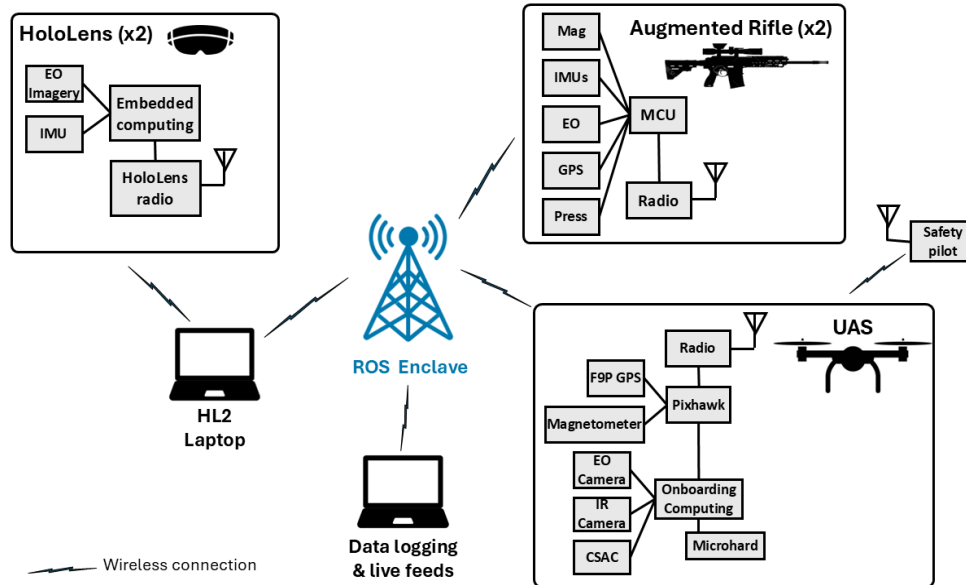


Figure 4. Network architecture schematic. Notes: EO = electro-optical; IMU = inertial movement unit.

3.4.1 HoloLens

For this data collection, the HoloLens (Microsoft Corporation; Redmond, Washington), an augmented reality (AR) headset with multimodal data recording capabilities, was used. This system is primarily designed to be used indoors. However, in current work, these headsets were used outdoors, which presented multiple challenges. Critically, the outdoor environment often offers less inherent structure (e.g., walls), which can be used by the headset as cues to create a map of the surroundings and track the user's position in relation to the origin point. Parallel Signal Integration for Enhanced Collaborative Human Observation (PSI-ECHO), software developed by DCS Corporation, was used to mitigate this challenge.

PSI-ECHO combines the multimodal integration of data streams from each headset (e.g., video and eye tracking) enabled by Microsoft's Platform for Situated Intelligence (PSI) with multiuser colocation abilities created through a shared spatial anchor across headsets. Critically, this allows time-synced multimodal data from multiple headsets to be collected from a shared coordinate frame. To build the shared coordinate frame, a common spatial anchor was set. Afterward, each Soldier confirmed that they saw the spatial anchor (represented by a holographic sphere) at the same location in the real world. This procedure was repeated for each data collection to ensure calibration was successful.

For each run, the HoloLens recorded POV video data, eye tracking, head tracking (rotation and location), and gesture detection. Audio was not recorded, because pilot data showed that its inclusion led to an increase in the rate of dropped samples

and system crashes. Although PSI-ECHO allows for multiuser data to be streamed and visualized in real-time using PSI Studio, for this collection, data was recorded and saved on the HoloLens HMDs to prevent dropped samples. Pilot testing was done to determine the optimum resolution and frame rate for recording video data. Ultimately, a frame resolution of 1920×1080 with a sampling rate of 15 Hz was determined to provide the fidelity required to run the computer vision algorithms without a large number of samples being dropped.

3.4.2 Custom Instrumented Rifles

The instrumented rifle (Figure 5) utilizes the following hardware components: a Tactacam 5.0 (Tactacam; Decorah, Iowa) camera for video data and a rigidly mounted sensor suite comprising an IMU, GPS receiver, and magnetometer. Custom software, including modified device drivers for the IMU and camera and associated communication protocols, supports these functions. Hardware was mounted on to the rifle's Picatinny rails using 3D-printed acrylonitrile butadiene styrene plastic components. A Linux device driver was used for the webcam to capture video data. A variety of resolutions were available with this camera. For the current data collection, an image resolution of 1280×720 was used with a sampling rate of 15 Hz. All data from the rifles were included in one ROS bag. Trigger pull timestamps were used to determine when targets were engaged. The other data streams captured on the rifles were stored but not analyzed or reported in this report.



Figure 5. Instrumented rifle with video, IMU, GPS, and magnetometer.

3.5 Inferred Shared Coordinate Space

Data from the HoloLens was wrangled using the SpatioTemporal Analysis and Inference Toolkit (STAIT) developed by ARL and DCS Corporation.* STAIT aligned multimodal samples (e.g., head tracking, eye tracking, and video) in time, using the shared coordinate system of the HoloLens, generated by the spatial anchors. Grounding DINO was then used to identify targets (red boxes) in each frame of video from each Soldier. The head location of each Soldier was plotted, and the yaw orientation, the left or right rotation, was used to determine their facing at the time of each frame. When a target was detected in a frame, the pixels of the frame were scaled to a 60°-wide field of view, set by the Soldier's head location and orientation at that moment. The x-coordinate of the center of a target's bounding box was used to set the target location in the left-right direction relative to the Soldier.

If we know where the Soldier is in space and have an approximation of the target's location to the left and right of the Soldier using the x-coordinate of the target bounding box in the POV video frame, our next step is to approximate the depth of the target, relative to the Soldier. To do this, the following formula was used, which takes the forward x,y,z unit vector of the Soldier's head orientation and calculates pitch (i.e., how far the head was tilted up or down).

$$pitch = \arcsin\left(\frac{z}{\sqrt{x^2 + y^2 + z^2}}\right) \quad (2)$$

Next, two distance modifiers were calculated. The first was *mod_{vertical}*. This modifier used the y-coordinate of a target's bounding box in the POV video frame to estimate depth (i.e., targets detected higher in frame are further away from the viewer). This was calculated with the bounding box's center y-coordinate and two additional constants.

$$mod_{vertical} = 0.2 + (box_y \times 1.3). \quad (3)$$

The second modifier, *mod_{pitch}*, accounts for the pitch of the Soldier's head, with the assumption that if a Soldier's head is angled upward, it is likely that the target in view is further away than if the Soldier's head was angled down to view it.

$$mod_{pitch} = \begin{cases} 1.0 + \frac{abs(pitch)}{45} & \text{if } pitch < 90 \\ .8 - \frac{pitch}{90.0} & \text{if } pitch > 90 \end{cases} \quad (4)$$

* SITCORE GitLab, <https://gitlab.sitcore.net/VisualSearch/sg-stait/>

Mod_{pitch} has two calculations depending on if the head is tilted up (i.e., $pitch > 90$) or down ($pitch < 90$). These modifiers combine with a preset constant to determine the target's distance from the Soldier in that frame.

$$Distance = \frac{7}{mod_{vertical} \times mod_{pitch}}. \quad (5)$$

For this initial pass, constants were set by trial and error, selecting those that yielded results closest to the ground truth based on visual inspection. To test this method, a confidence map was generated whereby the more frequently a given location was estimated to contain a target across frames, the higher the confidence score. Next, a 98% threshold was set so that only high confidence locations were compared to the ground truth locations of targets. Finally, an estimated target location was considered accurate if it fell within a given radius from a ground truth target. This accuracy calculation was performed iteratively, each time expanding the radius from the ground truth targets, over a range of 10 to 70 pixels.

3.6 Real UAS

Due to safety restrictions, the UAS could not search the MOUT site simultaneously with Soldiers. Instead, a Red Cat Teal Drones Teal 2 UAS was flown in between Soldier searches in a grid pattern over the site. The UAS collected position, heading, and gimbal angle via MISB-compliant Key-Length-Value metadata interleaved with the MPEG .ts video stream.

An Ultralytics YOLOv12 (large) model was utilized for object detection within the UAS video data. The base YOLOv12 object detection model was fine-tuned using custom annotated data from the test location itself. The custom data also included dataset augmentation, which added random scaling, rotation, translation, and flipping of the images (Figure 6).

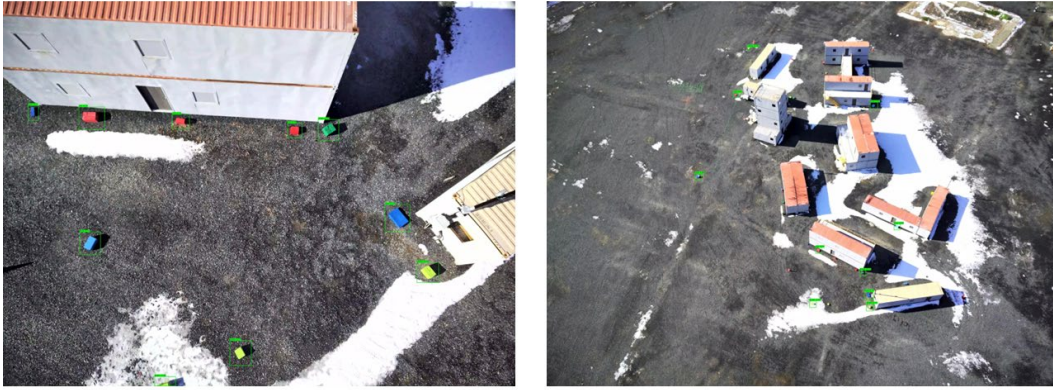


Figure 6. Example images from taken from the UAS at lower (left) and higher (right) altitudes. YOLO classification boundary boxes have been superimposed on images to show where objects (boxes) were detected.

4. Results

The following are results of the cognitive modeling, the simulated UAS flight, and the inferred coordinate space. The data collected from the rifle and real UAS are not presented here though some data points from the rifle (e.g., trigger pulls) were used to time lock HoloLens data streams for analysis.

4.1 Model

Averaged across two Soldiers, RAGNAROC predicted gaze location with an error of 28.10 dva. Using a saliency map (GBVS) alone to predict gaze location, an error of 29.86 dva was produced. Both models produced lower errors than the random item baseline that produced an estimation error of 39.83 dva (Figure 7).

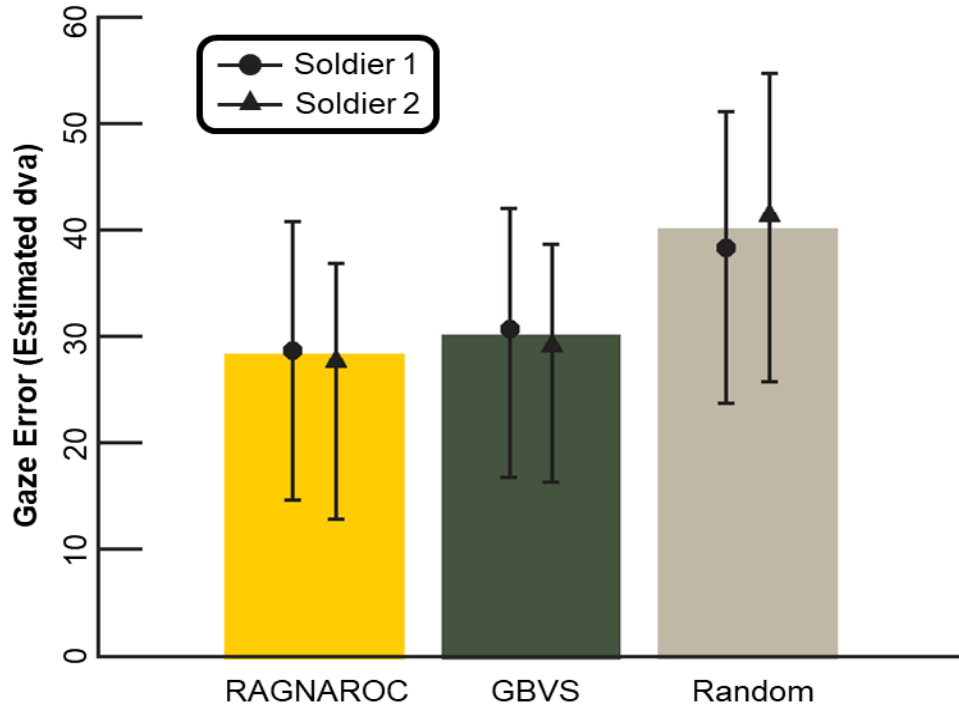


Figure 7. Gaze error produced by RAGNAROC, GBVS, and random selection. Error bars and markers indicate quantiles and mean values of individual Soldiers.

Cross-referencing trigger-pull timing and first-person video data, 11 valid target engagement events were identified. Three events were then excluded as the to-be-engaged target was the only visible target at the time of engagement. All of the remaining 8 events had at least 20 video frames (~ 1.33 s) before target engagement where both the to-be-engaged and control target objects were visible. When averaged across this period, attentional priority at the to-be-engaged target ($M = 2.16$) was higher than that at the control target ($M = 1.38$). Prioritization of the to-be-engaged target emerged at approximately 0.5 s pre-engagement (Figure 8). Z-scores are reported because attentional priority generated by RAGNAROC had an arbitrary unit.

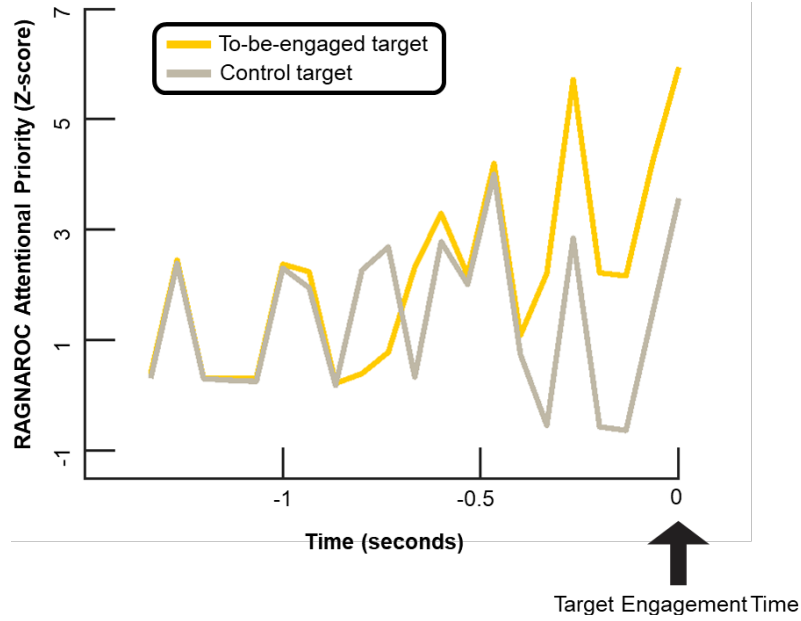


Figure 8. Attentional priority (Z-scores) in period before target engagement via rifle trigger pull. Golden line represents priority at the to-be-engaged target's location, and tan line represents priority at the location of a control target that was also visible

4.2 Simulated UAS

Across the run, human Soldiers covered 55.92% of on-ground area in the absence of simulated UAS support. Coverage was improved to 79.30% with support from independent UAS. Adaptive UAS, which adjusted flight speed through opportunistic sensing, led to the largest improvement where squad-level coverage reached 100% (Figure 9).

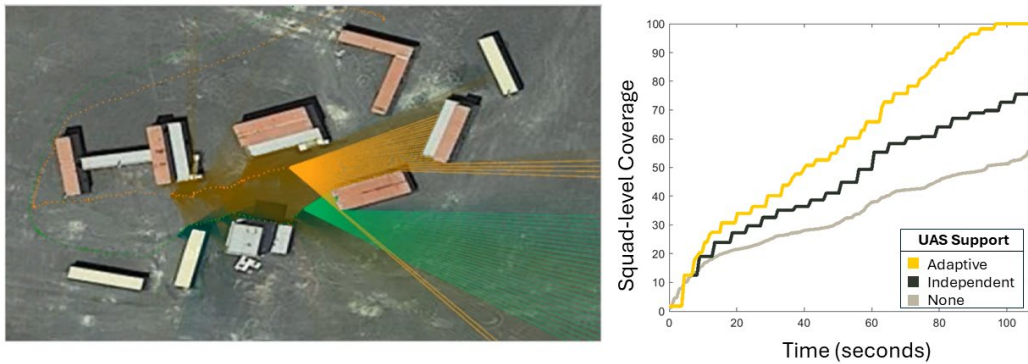


Figure 9. Squad-level coverage of the environment. Left: Visualization of projected field of regard for both Soldiers in orange and green, respectively. Dotted lines denote previous locations of either Soldier. Field of regard is plotted as a 70° wedge of rays, originating from the Soldier's head location and rotated using the Soldier's head orientation. Rays terminate at objects (e.g., buildings, defined by an edge detection algorithm run on the map image) and increase in transparency over time. Right: Squad-level coverage of the environment plotted across time to demonstrate the effect of incorporating UAS support with independent vs. adaptive programming of flight speed.

4.3 Inferred Shared Coordinate Space

Figure 10 illustrates the confidence map generated in estimated target locations in the MOUT site. After applying thresholding, the model was assessed on both recall (how many objects were detected) and precision (of the objects detected, how many were targets). In terms of recall, less than 10% of targets were detected at the smallest radius tested (10 pixels), with performance plateauing at 80% with a search radius of 60 pixels. Precision also hit peak performance (20%) with a search radius of 60 pixels.

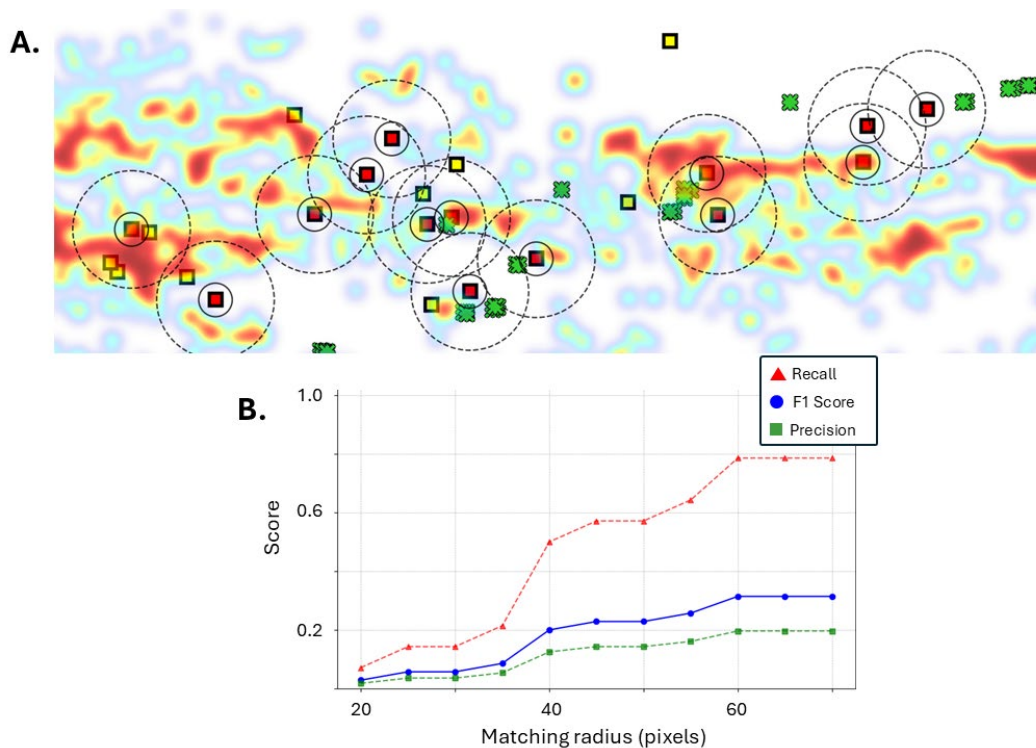


Figure 10. Results of inferred shared coordinate space. A. Confidence heat map. Warmer colors denote areas that were frequently predicted to contain targets across video frames. Red squares are ground truth locations of targets (red boxes). Yellow squares are ground truth locations of non-targets (yellow). Note that other non-targets were present in the environment, but locations were not stored. Green squares denote locations of estimated targets by placing a green square directly in front of the Soldier's head using the head pitch to estimate depth at the time of trigger pulls. Solid circles around target denote the smallest radius used in the radial sweep (20 pixels), and dotted circles represent the largest radius used (70 pixels). **B. Precision, Recall, and F1 score.**

5. Discussion

The goal of the current work was to demonstrate how leveraging cognitive science in conjunction with innovative hardware and software solutions can move the current state of the art toward a reality where SA of the environment is shared across

humans and autonomous agents for enhanced lethality. To this end, a dual-pronged approach was used to build a representation of the environment, as understood by humans and machines, that was shareable across members of the heterogeneous squad. To characterize the human's internalized SA, a cognitive model was applied to a real-world, mission-relevant visual search task performed by Soldiers at an outdoor MOUT site. Because this data collection was a proof of concept, no statistical analyses were performed on the model validation metrics. However, qualitatively, the model showed enough evidence of capturing human signal that further research in this area is warranted. To map that awareness into a shared coordinate space, methodologies were developed to colocate the Soldiers in the world with the objects (targets) that surrounded them. Finally, a joint search was simulated using the Soldiers' search pattern data and pairing them, in simulation, either with a UAS flying independently or adaptively altering its flight speed based on sensed human information. Although this was a preliminary instantiation of such an HMI mission, the results demonstrate the potential performance enhancements available when autonomous systems are informed by the human's awareness of the environment.

5.1 Outstanding Challenges and Future Solutions

The results of this demonstration indicate that this area of research could be impactful in improving the effectiveness of dismounted heterogeneous units. However, there are several outstanding challenges that need to be addressed. In this report, each challenge identified is discussed, including how the current work made progress towards this end, what gaps remain, and how future work might potentially bridge those gaps.

5.1.1 Translating Retinotopic 2D AM Activation to a 3D World Map

The original intent of this work was to build a squad-level SA map by translating RAGNAROC's estimated AM activation into a 3D coordinate space that could inform the UAS flight path. Although the model was successfully run on POV video from each Soldier, projecting the 2D activation of each video frame to the 3D world proved to be challenging without ground truth knowledge of the environment. Specifically, not knowing the distance between the Soldier and objects around them made this 3D projection difficult. As a placeholder, the simulated UAS was fed field-of-regard information from each Soldier. In other words, a binary was instituted whereby all locations within the Soldier's field of regard were coded as perceived. Moreover, the field of regard was represented in a 2D bird's-eye view map for the environment, using building edges as occluders to a Soldier's line of sight. This approach sacrifices detail about information perceived

in the vertical dimension (e.g., windows on the upper floors of buildings), which would be necessary to a valuable shared representation of SA in a real-mission scenario. Therefore, the outstanding goal remains to build a 3D, SA map that not only represents the information that was perceptible to the human but weights that information according to estimated attentional prioritization because not all locations in the visual field are perceived to the same degree (Mack 2003; Kastner and Pinsky 2004; McMains and Somers 2004).

To address this, future work will move this scenario into VR where ground truth of the environment is known and controlled. With precise measurement of the distance between the viewer and objects in their environment, more accurate projections can be made of the 2D AM activation onto the 3D world. Then a simulated drone can be flown, guided by this 3D human attentional landscape, allowing us to test whether access to such information improves human–UAS joint search over and above simply providing field-of-regard information.

5.1.2 Developing and Testing Alternative Adaptive UAS Flight Algorithms

In the current work, a relatively simple algorithm was used to adapt UAS flight patterns to human data. Namely, the environment was split into a grid, and the time in which a UAS spent processing a given section was inversely proportional to the amount of time that section was collectively in view of the Soldiers. However, this is not the only way to inform flight patterns.

Future work will develop alternative adaptive algorithms to test how access to human data could be used to decrease power consumption and increase flight path efficiency. Towards that end, development is underway to adapt the pathfinder algorithm that is designed to find a path through the environment that limits the user's time within an enemy's line of sight (Strube and Oates 2023). In the original experiment by Strube and Oates (2023), enemy locations are plotted on a map and the areas that are blocked from their line of sight are marked. The pathfinder algorithm then charts a course for participants to move through the environment from a start to finish location, minimizing the amount of time spent in the enemy's line of sight. For the current work's purposes, this algorithm will be adapted so that instead of avoiding an enemy's line of sight, the path avoids locations already heavily attended to by the human squad. Developing this and other competing algorithms will be essential in testing how the human context can be used most effectively.

5.1.3 Raising the Signal-to-Noise Ratio in the Inferred Shared Coordinate System

Using head angle and location in combination with a computer vision object detection algorithm run on POV video data showed promise in being able to plot the Soldiers' positions in the same coordinate system with objects in the environment. Building this capability is an essential component of being able to translate AM activation to a 3D coordinate space in a real-world data collection (see Section 5.1.1). However, in its current form, this methodology captures too much noise for this to be a viable option. The source of much of the noise is due to spurious target classifications by Grounding DINO and therefore could be improved by other, potentially better-trained, computer vision algorithms. However, incorporating other streams of data from humans might also improve performance.

Currently, when a target is detected in a given frame of the POV video, the depth of that target is determined by the pitch of the Soldier's head. This is a proxy metric because objects that are further away are more likely to elicit a positive pitch, where the head is tilted up to view the object compared to when said object is closer to the viewer. Although this is a reasonable heuristic, it does not account for the possibility that a target may be detected in a frame, without being the actual object of fixation. In other words, the target in view might not be the object the viewer is currently looking at and therefore orienting their head to accommodate. Therefore, although head pitch may help set a range for possible depth of the objects in a frame, it should be more tightly correlated with the depth of the object on which the viewer is currently fixating gaze. Consequently, incorporating eye tracking information may prove useful in specifying object location estimations. One such approach that was developed using this data collection used the eye unit vector along with an assumed height of the Soldier to triangulate the approximate location of fixated targets (Saravana and Cohen Hoffing 2025). However, a key limitation of this work is the assumption that all identified targets are at ground level. This puts constraints on its usability in the field and highlights the possible need for incorporating ML solutions for general scene understanding. For instance, it may be necessary to not only detect targets but also identify other objects or landmarks in the environment and understand their relationships (e.g., the target detected in frame is on top of a building).

5.1.4 Using Opportunistic Sensing to Improve UAS Target Classification

Ultimately, this work pushes toward a future where UAS and humans could search the environment in parallel, informed by each other to improve overall SA and enhance lethality. Toward that end, a computer vision algorithm was run on the

video data obtained from the UAS to detect target and nontarget objects (i.e., boxes of all colors). A hand-coded ground truth was not established so model accuracy was not calculated or reported. However, through visual inspection of the model's object classifications, it appeared that the model performed better at lower altitudes when the boxes were bigger in frame. At higher altitudes, model accuracy declined, requiring more iterations of the semi-supervised training regime to maintain performance. Flying the UAS at higher altitudes is advantageous because it allows for a larger field of view or search area. Therefore, being able to detect targets at distance is critical for enhancing overall unit coverage of the environment. In the current method, the computer vision algorithm was retrained in part by presenting hand-labeled images of boxes with different rotations and at different scales. Inducing this sort of variation of images has been shown to increase robustness in classification and detection (Zhao et al. 2025). Establishing human-labeled ground truth images is the burdensome part of this process.

To mitigate this burden, future work will focus on using opportunistically sensed ground truth from the Soldier's POV video data. Specifically, using the time stamps of the rifle's trigger pulls will allow us to select human-verified frames in which targets were in view (i.e., because the Soldiers confirmed that a target was present by shooting at it). These frames, labeled by the human without additional burden (because shooting targets was part of their primary task), can then be used to help train target detection models run on the UAS video data, providing images of the same objects from different viewpoints and at different scales. This future work will be able to quantify whether opportunistically sensed data reduces training time to improved target detection accuracy in a paired robotic system.

5.1.5 Improving Hardware and Software Support for Real-World Multi-Human Team Data Collections

To carry out the current work, state-of-the-art network capabilities were built in-house including novel hardware and software solutions. Namely, the HoloLens headsets and the augmented rifles were adapted not only to support this data collection but intended to serve as a foundation for future experimental work as the work under DIESL scales up to include multi-human, multi-agent units. Toward that end, several modifications are currently underway as a result of the lessons learned from this work.

For the rifle, future modifications include making sensors more compact and encapsulated in the rifle itself, rather than being mounted on top. This will not only help make using the rifle more natural for the Soldier (allowing us to capture a closer approximation to behavior on the battlefield), but it will also help to ruggedize the technology because weather conditions in outdoor experimentation

proved to be variable. In addition, further work is being done to decrease the amount of power used by the sensor suite. This will allow for longer data collections and expand our potential breadth of research (e.g., studying fatigue or training effects). Finally, work is being done both on the HoloLens side as well as the rifle side to mitigate data loss. For the rifle, this involves building solutions so that data can still be stored on a device even when network connections are dropped. For the HoloLens, this involves building remote shut down capabilities so that data can be salvaged in the event of a system crash.

Finally, an essential component (and significant challenge) to outdoor, real-world experimentation is ground truth knowledge of the environment. Toward that end, in future work, we will build a virtual model of the experimental space (e.g., the MOUT site) and place targets within the virtual space using a game engine such as Unity. These locations could be displayed via the HoloLens to the experimenter during set up through AR holograms to instruct them on where to put real-world stimuli (e.g., target boxes). This way, in addition to having the shared spatial anchor using PSI-ECHO to plot Soldier's positions, object locations within that same coordinate space would also be available. This would provide a solid ground truth by which to test algorithms aimed at inferring such a coordinate space.

6. Conclusion

Although advancements in autonomous robotics systems offer huge potential in enhancing Warfighter effectiveness, these enhancements are limited by the degree to which machines can be integrated into units with humans. A key component of this integration is a shared SA or internal representation of the surrounding environment—not only understanding what information the human has gathered but also knowing what information was likely missed and how robotic systems might respond adaptively to fill those gaps. The current work served as a demonstration of how cognitive theory, developed in the laboratory, can be leveraged toward this end and what experimental and technical hurdles currently remain.

7. References

- Endsley MR, Garland DJ. Theoretical underpinnings of situation awareness: a critical review. *Situation awareness analysis and measurement*. Lawrence Erlbaum Associates; c2000. p. 3–21.
- Harel J, Koch C, Perona P. Graph-based visual saliency. *Advances in Neural Information Processing Systems 19: Proceedings of the 2006 Conference*; 2007. MIT Press.
- Kastner S, Pinsk MA. Visual attention as a multilevel selection process. *Cogn Affect Behav Neurosci*. 2004;4(4):483–500.
- Lance BJ et al. Minimizing data requirements for soldier-interactive AI/ML applications through opportunistic sensing. *Artificial Intelligence and Machine Learning for Multi-Domain Operations Applications II*. SPIE. 2020; Online. 11413.
- Li Z. A saliency map in primary visual cortex. *Trends Cogn Sci*. 2002;6(1):9–16.
- Liu S et al. Grounding DINO: marrying DINO with grounded pre-training for open-set object detection. Paper presented at: *European Conference on Computer Vision*; 2024 Sep 29–Oct 4; Milan, Italy. p. 38–55.
- Lucas BD, Kanade T. An iterative image registration technique with an application to stereo vision. *Proceedings of the 7th International Joint Conference on Artificial Intelligence (IJCAI '81)*; Vancouver, Canada. c1981; p. 121–130.
- Mack A. Inattention blindness: looking without seeing. *Curr Dir Psychol Sci*. 2003;12(5):180–184.
- McMains SA, Somers DC. Multiple spotlights of attentional selection in human visual cortex. *Neuron*. 2004;42:677–686.
- Saravana S, Cohen Hoffing R. Multimodal 3D gaze-object localization for shared human-machine situational awareness in tactical environments. *DEVCOM Army Research Laboratory (US)*; 2025. Report No.: ARL-TR-10183.
- Shi J. Good features to track. Paper presented at: *IEEE Conference on Computer Vision and Pattern Recognition (CVPR94)*; 1994 June; Seattle, Washington.
- Strube E, Oates K. Seeing through the enemy's eyes. *Proceedings Volume 12538, Artificial Intelligence and Machine Learning for Multi-Domain Operations Applications V*; 2023 June; Orlando, Florida. SPIE Defense + Commercial Sensing. <https://doi.org/10.1117/12.2659166>

- Szeliski R. Computer vision: algorithms and applications. Springer Nature; 2022.
- Tam J, Callahan-Flintoft C. Neural computational model predicts attentional dynamics during immersive search in virtual reality. bioRxiv. 2025 Oct 7. <https://doi.org/10.1101/2025.10.07.680957>
- Tian Y, Ye Q, Doermann D. Attention-centric real-time object detectors. arXiv. 2025 Feb. arXiv:2502.12524. <https://doi.org/10.48550/arXiv.2502.12524>
- Tomasi C, Kanade T. Detection and tracking of point. Int J Comput Vis. 1991;9:137–154.
- Wang C, Yang X, Xi X, Nie S, Dong P. Introduction to LiDAR remote sensing. CRC Press; 2024.
- Wyble B et al. Understanding visual attention with RAGNAROC: a reflexive attention gradient through neural attractor competition. Psychol Rev. 2020;127(6):1163.
- Zhao Z, Bakar EB, Razak NBA, Akhtar MN. Progressive optimization of EfficientNetV2 for classification based on images of corroded objects. Comput Appl Math. 2025;44(8):419.

List of Symbols, Abbreviations, and Acronyms

2/3D	Two-/Three-Dimensional
AI	Artificial Intelligence
AM	Attention Map
AR	Augmented Reality
ARL	Army Research Laboratory
DEVCOM	U.S. Army Combat Capabilities Development Command
DIESL	Distributed Information for Enhanced Squad Lethality
dva	Degrees of Visual Angle
GBVS	Graph-Based Visual Saliency
GPS	Global Positioning System
HMD	Head-Mounted Display
HMI	Human–Machine Integration
IG	Inhibitory Gating
IMU	Inertial Motion Unit
LiDAR	Light Detection and Ranging
ML	Machine Learning
MOUT	Military Operations on Urbanized Terrain
POV	Point of View
PSI	Platform for Situated Intelligence
PSI-ECHO	Parallel Signal Integration for Enhanced Collaborative Human Observation
R2C2	Robotics Research Collaboration Campus
RAGNAROC	Reflexive Attention Gradient through Neural Attractor Competition
ROS	Robotic Operating System
SA	Situational Awareness

STAIT	Spatiotemporal Analysis and Inference Toolkit
UAS	Unmanned Aerial System
VR	Virtual Reality
YOLO	You Only Look Once

1 DEFENSE TECHNICAL
(PDF) INFORMATION CTR
DTIC OCA

1 DEVCOM ARL
(PDF) FCDD RLB CI
TECH LIB