

Model-Driven Requirements for Humans-on-the-Loop Multi-UAV Missions

Ankit Agrawal and Jane Cleland-Huang
Dept. of Computer Science and Engineering
University of Notre Dame
Notre Dame, IN, USA
aagrawa2@nd.edu, JaneHuang@nd.edu

Jan-Philipp Steghöfer
Dept. of Computer Science and Engineering
Chalmers | University of Gothenburg
Gothenburg, Sweden
jan-philipp.steghofer@cse.gu.se

Abstract—The use of semi-autonomous Unmanned Aerial Vehicles (UAVs or drones) to support emergency response scenarios, such as fire surveillance and search-and-rescue, has the potential for huge societal benefits. Onboard sensors and artificial intelligence (AI) allow these UAVs to operate autonomously in the environment. However, human intelligence and domain expertise are crucial in planning and guiding UAVs to accomplish the mission. Therefore, humans and multiple UAVs need to collaborate as a team to conduct a time-critical mission successfully. We propose a meta-model to describe interactions among the human operators and the autonomous swarm of UAVs. The meta-model also provides a language to describe the roles of UAVs and humans and the autonomous decisions. We complement the meta-model with a template of requirements elicitation questions to derive models for specific missions. We also identify common scenarios where humans should collaborate with UAVs to augment the autonomy of the UAVs. We introduce the meta-model and the requirements elicitation process with examples drawn from a search-and-rescue mission in which multiple UAVs collaborate with humans to respond to the emergency. We then apply it to a second scenario in which UAVs support first responders in fighting a structural fire. Our results show that the meta-model and the template of questions support the modeling of the human-on-the-loop human interactions for these complex missions, suggesting that it is a useful tool for modeling the human-on-the-loop interactions for multi-UAVs missions.

Index Terms—Human Multi-Agent Collaboration, Requirements Elicitation, Autonomous Agents

I. INTRODUCTION

The deployment of a swarm of Unmanned-Aerial Vehicles (UAVs) to support human first responders in emergencies such as river search-and-rescue, hazardous material sampling, and fire surveillance has earned significant attention due to advancements in the robotics and Artificial Intelligence (AI) domains [1], [2]. Advanced AI models can assist UAVs in performing tasks such as creating a 3D heat-map of a building, finding a drowning person in a river, and delivering a medical device, while robotics autonomy models enable UAVs to automatically plan their actions in a dynamic environment to achieve a task [3], [4]. However, despite these advances, the deployment of such systems remains challenging due to uncertainties in the outcome of the AI models [5], rapid changes in environmental conditions, and emerging requirements for how a swarm of autonomous UAVs can best support first responders during a mission.

The UAVs of next-generation emergency response systems will be capable of sensing, planning, reasoning, sharing, and acting to accomplish their tasks [6]. These UAVs will not require humans-in-the-loop to make all key decisions, but rather will make independent decisions with humans-on-the-loop setting goals and supervising the mission [7]. For example, in a multi-UAV river search-and-rescue mission, the autonomous UAV can detect a drowning person in the river utilizing the on-board AI vision models (*sensing*) and ask another UAV to schedule delivery of a flotation device to the victim's location (*planning* and *reasoning*). These UAVs collaborate to share (*sharing*) the victim's location and subsequently deliver the flotation device (*acting*). These intelligent UAVs also send the victim's location to emergency responders on the rescue-boat so that they can perform the physical rescue operation. Autonomous systems of such complexity demand humans and intelligent agents to collaborate as a human-agent team [8], [9].

A well-known issue in designing a system comprising humans and autonomous agents is to identify how they can collaborate and work together to achieve a common goal [10]. The challenges in human multi-agents collaboration include identifying when and how humans should adjust the autonomy levels of agents, identifying how autonomous agents should adapt and explain their current behavior to maintain humans' trust in them, and finally, identifying different ways to maintain situational awareness among humans and all autonomous agents. In this paper we propose a humans-on-the-loop solution in which humans maintain oversight while intelligent agents are empowered to autonomously make planning and enactment decisions. We first identify common interaction patterns in which humans collaborate with autonomous agents, and then leverage those patterns to construct a human interaction meta-model. In addition, we define a set of 'probing' questions which can be used to elicit, analyze, and ultimately specify requirements for human multi-UAV interactions in specific emergency response missions.

This paper makes three primary contributions. First it motivates the problem of human multi-agent interaction through examples drawn from a concrete mission scenario. Second, it provides a meta-model to describe human interactions with multiple agents, and finally it presents a set of requirements-related guiding questions for eliciting and then modeling specific instances of these human multi-agent interactions.

The paper is organized as follows: Section II presents examples of human multi-agent interactions drawn from the river-rescue scenario and section III presents an analysis of these interactions. Section IV introduces a human-on-the-loop meta-model for describing human multi-agent interactions. Section V then describe our process for eliciting requirements, mapping them to elements of the meta-model, and then specifying requirements by deriving instances of the meta-model for each identified human multi-agent interaction type. Section VI discusses an application of our work and finally sections VII, VIII, and IX discuss threats to validity, related work, and draw conclusions.

II. HUMAN-MULTI-UAV COLLABORATIONS

Several research groups have explored the application of UAVs for specific emergency scenarios such as surveying and assessing damage following an earthquake [11] or volcanic eruption [12], investigating maritime spills [13], delivering defibrillators [14], and mapping wildfires [15]. These applications all involve human operators interacting with UAVs in direct or indirect ways to plan routes, capture video, or to supervise varying degrees of autonomous UAV behavior – typically through the use of a graphical user interface (GUI). Researchers have described other forms of interactions [16], including haptic and voice interfaces [17], [18], but these are infrequently used in emergency response applications.

A. DroneResponse: A Case Environment

In this paper, we primarily draw examples from our *DroneResponse* system, which we are developing to enable multiple collaborating, semi-autonomous UAVs to support diverse emergency response missions such as fire surveillance, search-and-rescue, and environmental sampling [19], [20], [21]. Figure 1 depicts a river search-and-rescue use-case in which multiple UAVs are deployed to find a victim on the river and to potentially aid emergency responders in delivering a flotation device.

DroneResponse represents a socio-technical cyber-physical system (CPS) in which multiple humans and multiple semi-autonomous UAVs engage in a shared emergency response mission. UAVs are designed to make autonomous decisions based on their current goals, capabilities, and current knowledge. They build and maintain their knowledge of the mission through directly observing the environment (e.g., through use of their onboard sensors) and through receiving information from other UAVs, central control, and human operators [22]. UAVs then work to achieve their goals through enacting a series of tasks [23].

Humans interact with UAVs through various GUIs to create and view mission plans, monitor mission progress, assign permissions to UAVs, provide interactive guidance, and to maintain situational awareness. Bidirectional communication is crucial for enabling both humans and UAVs to complement each other’s capabilities during the mission. An example of human-UAV collaboration is depicted in Figure 2, which shows a UI developed for the *DroneResponse* system. In this

Use Case: River Search-and-Rescue

Pre-Conditions

Multiple UAVs are equipped with cameras and activated
Firefighters have marked area of river to be searched
DroneResponse is running and UAVs are displayed on map
One victim is in the search area

Main Success Scenario

1. Drones takeoff and commence the search.
2. DroneResponse tracks and displays the location and state of each UAV
3. The UAVs takeoff and execute their assigned search patterns
4. Each UAV processes imagery from its camera using a trained computer vision model.
5. One UAV (U1) finds and detects the victim with confidence greater than a predefined threshold.
6. The UAV raises an alert and switches to ‘active_tracking’ mode. The Drone Commander (DC) views the imagery and confirms that the victim has been found.
7. The DC confirms ‘active tracking’
8. The DC notifies the Incident commander who confirms the sighting & directs human responders in their boat to rescue the victim.
9. The DC confirms that UAV (U1) has sufficient battery to continue active-tracking. He/she recalls all other drones to their home-base.
10. Human responders arrive at the scene.
11. The DC recalls all UAVs to their home base.

Fig. 1: A partial use case description of the DroneResponse River search-and-rescue scenario.

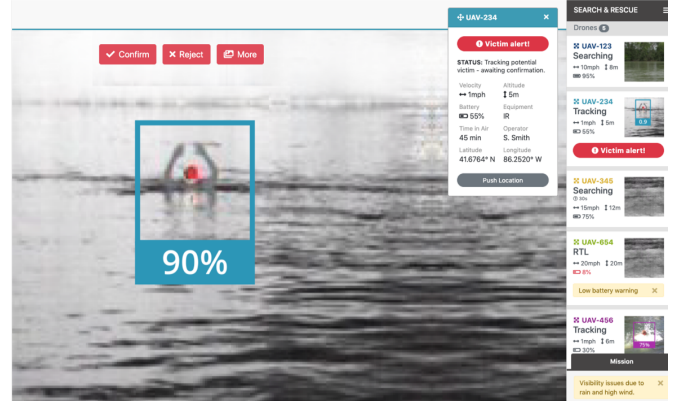


Fig. 2: A human-UAV interaction point in which a UAV has detected a candidate victim and requested human confirmation.

example, the UAV has detected a candidate victim in the water, autonomously started tracking the victim, while simultaneously requesting confirmation from the human incident commander that the detected object is actually the victim.

B. Human-UAV Interactions

DroneResponse is being developed in close collaboration with emergency responders through engagement in a series of brainstorming activities, interviews, participatory design sessions, and early field-tests [19], [20], [24]. The following concrete examples of human-UAV interactions, taken from the river search-and-rescue example, were identified as part of this collaborative design process. We use these examples throughout the remainder of the paper to motivate and contextualize our modeling activities.

Scenario S1 – Planning a rescue strategy: When a UAV identifies a potential victim in the river, the victim’s coordinates are sent to the mobile rescue unit. However, the UAV must also decide whether to request delivery of a flotation device by a suitably equipped UAV or whether it is sufficient to simply continue streaming imagery of the victim until human rescuers arrive. The UAV makes this decision by estimating the arrival time of the rescue boat versus the time to deliver a flotation device. However, humans can contribute additional information to the decision – for example, by modifying the expected arrival time of the rescue boat, or by inspecting the streamed imagery and determining whether the victim would be able to receive the flotation device if it were delivered (e.g., the victim is conscious and not obscured by overhead branches) and is in need of the device (e.g., not having a safe waiting position on a rock or tree branch). This is an example of a *bidirectional exchange of knowledge* between multiple humans and multiple UAVs, where the first UAV shares the victim’s coordinates and streams imagery, humans on the boat estimate their ETA and if necessary update the UAV’s situational awareness, the incident commander decides whether a flotation device could be used effectively if delivered on time, and if needed, a second UAV performs the delivery. The scenario illustrates many aspects of human-agent collaboration including *knowledge sharing* and *human intervention*.

Scenario S2 – Sharing environmental information: In river search-and-rescue missions, victims tend to get trapped in ‘strainers’ (i.e., obstruction points) or tangled in tree roots on outer banks. These areas require closer inspection. While UAVs have onboard vision and will attempt to identify ‘hotspots’, human responders can directly provide this information to multiple UAVs based on their observation of the scene. This enables UAVs to collaboratively adapt their flight plan so that they prioritize specific search areas, or adjust their flight patterns to reduce speed or fly at lower altitudes in order to render higher-resolution images of priority search areas. This interaction scenario is similar to the previous one, except that it is primarily uni-directional with information passed from humans to UAVs.

Scenario S3 – Victim confirmation: The UAV’s AI model uses its onboard computer vision to detect potential victims. When the confidence level surpasses a given threshold, the UAV will autonomously switch to tracking mode and broadcast this information to all other UAVs. If the UAV autonomy level is low, it requests human confirmation of the victim sighting before it starts tracking. Human feedback is sent to the UAV and propagated across all other UAVs. In this scenario the *UAV elicits help from the human* and the human responds by confirming or refuting the UAV’s belief that it has sighted a victim or by suggesting additional actions. For example, if the detected object is partially obscured, the human might ask the UAV to collect additional imagery from multiple altitudes and angles.

Scenario S4 – Support for UAV coordination: In an extension to the previous scenario, multiple UAVs might

simultaneously detect a victim. They must then use onboard computer vision and their own estimated coordinates of the detected object to determine whether they have detected the same object and to plan a coordinated response. However, this determination may be more complicated in poor visibility environments with weak satellite signals and low geolocation accuracy (e.g., in canyons). Human responders may need to intervene in the UAV’s planning process by helping determine whether the sighted objects are valid and unique, and if necessary selecting the most appropriate UAV for the tracking task. This is an example in which the human *intervenes in the UAV’s autonomy* and potentially provides *direct commands*, assigning a specific UAV to the task.

Scenario S5 – Prohibiting normal behavior: Most UAVs come with built-in safety features so that they autonomously land-in-place or return to launch (RTL) when their battery becomes low or a malfunction is detected. In the case of a low battery, the DroneResponse system initially raises a low-battery alert in the UI, and eventually initiates the RTL command. A human responder might modify the UAV’s permissions and *prohibit the UAV from transitioning to RTL* if the UAV is conducting a critical task. An example, that arose from discussions with the Navy, was the use of floating drones for man-overboard scenarios. If a UAV found a victim, and no other UAV or human rescue unit were in the vicinity, the RTL feature would be deactivated automatically. This meant that when batteries lost power, the UAV would land in the water and serve as a search beacon. However, for many reasons, a human might wish to override the default deactivation of the RTL, thereby reactivating the UAV’s RTL autonomy.

These motivating examples provide the foundation for our discussion of human-on-the-loop collaboration patterns.

III. ANALYSIS OF COLLABORATION ACTIONS

Agents within a human-on-the-loop (HotL) system are empowered to execute tasks independently with humans serving in a purely supervisory role [25]. However, as our previous examples have shown, humans and agents continually share information in order to maintain bidirectional situational awareness and to work collaboratively towards achieving mission goals. Agents report on their status (e.g., remaining battery levels, GPS coordinates, and altitude), and they explain their current plans, actions, and autonomous decisions whenever requested by humans. Humans can directly intervene in the agents’ behavior by providing additional information about the environment, and agents can then leverage this information to make more informed decisions. Humans also respond to direct requests for feedback – for example, to confirm a victim sighting as previously discussed. They can also provide direct commands (e.g., RTL or stop tracking), or can explicitly modify an agent’s permissions in order to enhance or constrain the agent’s autonomous behavior. These types of interactions are depicted in Figure 3.

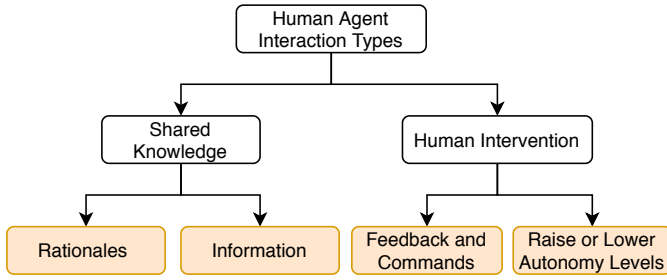


Fig. 3: Humans and agents collaborate through shared knowledge and through human interventions in the agents’ autonomy.

A. Situational Awareness

Situational Awareness (SA) is the ability of the user to perceive the environment (Level-1 SA), to understand the reasoning behind the current state of the environment (Level-2 SA), and finally, to project how the situation could evolve in the future (Level-3 SA) [26]. Humans acquire knowledge of the situation from diverse sources such as their physical interactions with the agents (e.g., visual observations and sounds), observations of the current weather, radio communication with on-scene first responders, and finally through information shared through the systems’ GUI. Humans combine knowledge from all of these sources to create a mental model of the current status of the mission. At the same time, autonomous agents, such as UAVs, develop their own situational awareness using their onboard sensors and through collating information shared by other autonomous agents and by humans. Both humans and autonomous agents then use their shared knowledge of the environment to formulate and enact plans to collaboratively achieve their mission goals.

In a HotL environment, agents make many autonomous decisions; however, in order for humans to supervise the mission and to maintain full situational awareness, the agents must explain their behavior when requested by a human. The explanation should include key **information** (i.e., the agent’s situational awareness at the time the decision was made), the autonomous decision (e.g., switch modes, change altitude etc), and a human understandable rationale for the decision. Providing **rationales for all decisions and subsequent behavior** is therefore critical in order for humans to achieve situational awareness. If the human were to disagree with the decision and the logic of the supporting rationale, then they could monitor the agents more closely, temporarily lower their autonomy levels, or make longer-term adjustments (e.g., retraining a computer vision model) for future missions.

B. Human Intervention

At times, humans may need to intervene in the autonomy of an agent in order to influence and improve the outcome of the joint mission. They can do so in several different ways. Previous studies[27], [28] demonstrate that a **feedback loop** can help agents to improve their future performance by fine-tuning algorithmic parameters that drive the agent’s autonomy.

For example, feedback on a candidate victim detected by the computer vision model, could be used to retrain the model or refine its configuration parameters, thereby potentially reducing false positives or false negatives. In addition, users can initiate commands to immediately enact changes in the behavior of the UAV. For example, a human could directly command a UAV to fly to a specific waypoint to checkout a report received on social media.

Finally, the human may choose to **raise or lower autonomy levels** of the agent. Autonomy levels, defined as the extent of an agent’s independence while acting autonomously in the environment, can be expressed through role assignments or through specific permissions within a role. For example, a UAV that is permitted to track a victim without first obtaining human confirmation has a higher autonomy level than one which needs explicit human confirmation before tracking. Humans tend to establish autonomy levels based on their trust in the agent’s capabilities. For example, a UAV exhibiting high degrees of accuracy in the way it classifies an object increases human trust, and as a result, the human might grant the UAV additional permissions. On the other hand, the human operator might revoke permissions, thereby lowering autonomy levels, if the UAV were operating in weather conditions for which the computer vision model had not been appropriately trained and for which accuracy was expected to be lower than normal.

IV. META-MODEL FOR HUMAN-UAV INTERACTIONS

We constructed a meta-model to define the vocabulary of the domain of human multi-agent interactions. The meta-model includes domain-specific concepts and establishes rules for how those types of concepts are associated with one another. This allows us to express specific instances of human multi-agent interaction in conceptual models and reuse the concepts we identified to express how humans and multiple agents will interact with each other in specific scenarios.

TABLE I: Additional Use-Cases from which human multi-UAV interaction patterns were identified and analyzed

ID	Use Cases	Engaged Stakeholders
UC1	River Search & Rescue	South Bend Firefighters
UC2	Defibrillator Delivery	DeLive, Cardiac Science
UC3	Traffic Accident surveillance	South Bend Firefighters
UC4	Water Sampling	Environmental Scientists
UC5	Man overboard	US Navy

The elements of the meta-model were (cf. Fig. 4) derived from our analysis of human multi-UAV interactions in the river-rescue scenarios and also from additional scenarios summarized in Table I. The meta-model depicts frequently occurring concept types and their associations, and was designed iteratively through multiple refinements in which we recursively validate the model against the specific scenarios described in Section II. Our meta-model includes the following elements:

A **Role** defines the complex behaviors that agents perform autonomously. Complex behaviors of a UAV include takeoff, search, track, deliver, and RTL.

An `AutonomousDecision` entity uses algorithms that leverage *Information* in the `KnowledgeBase` to make decisions. The complex behaviour of a `Role` is defined through one or several such decisions. For example, there are many cases in which a single agent must serve as a *leader*, responsible for coordinating behavior of its *followers*. During a leader election, an `AutonomousDecision` entity could select a new leader from the set of followers, thereby enabling the system to switch leaders without the need for human intervention. Upon making a decision, an `AutonomousDecision` entity generates output `Information` including a rationale for its decision, which could later be used to generate a human-readable explanation.

Entities of type `Permission` are used by `AutonomousDecisions` to decide if the agents are allowed to make a specific decision. For example, an `AutonomousDecision` entity checks whether the human responders have allowed the system to automatically select a replacement if needed during a victim tracking activity. Roles are associated with a set of permissions defining the allowed behaviors of the agent which can be modified at run-time.

A `KnowledgeBase` entity contains current environmental information as well as information about the state of a single agent or multiple agents. An `AutonomousDecision` entity uses the `Information` stored in the `KnowledgeBase` in decision making. A human can use the information in the `KnowledgeBase` entity to gain situational awareness of the mission.

Entities of type `HumanInteraction` allow humans to intervene in the autonomy of the agents or to share their explicit knowledge of the environment. The three entity types `ProvidedInformation`, `ChangedPermission`, and `IssuedCommand` provide different ways for humans to interact with the system. The `ProvidedInformation` entity adds `Information` to the `KnowledgeBase` of the system to maintain the consistent knowledge among multiple agents. Humans can use interventions of type `ChangedPermission` to raise or lower the autonomy of an agent, or agents, based on their trust in the ability of the agents to make correct decisions within the current environment. Finally, an `IssuedCommand` entity allows humans to gain control over the autonomous behavior of the agents. For example, if a UAV loses communication with other UAVs in the mission and fails to deliver the flotation device when it is needed, a human can send a direct command that sets the current `Role` of the UAV to *deliver flotation device*.

It is noteworthy that neither humans nor agents are represented explicitly in our meta-model. The underlying implicit assumption is that roles are assigned to agents according to the capabilities of each UAV, that UAVs can assume new roles according to the state of the environment, constrained by permissions associated with their capabilities. Furthermore, humans and agents have access to one or several instances of the distributed `KnowledgeBase` which stores information acquired from the environment, multiple UAVs, and from

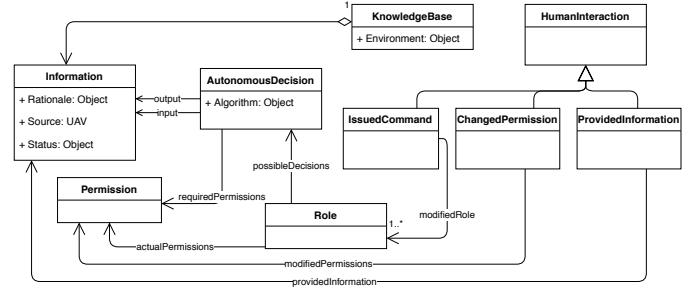


Fig. 4: Meta-model for human-on-the-loop interaction in a multi-agent mission.

humans. The reason for leaving these aspects implicit are that the domain of our model is human multi-UAV interaction and it is not relevant to the meta-model to specify which concrete UAV has assumed each specific role.

V. REQUIREMENTS MODELING

Human multi-agent interactions in the domain of emergency missions are impacted by factors such as uncertainty of the agents' knowledge, the degree of human trust in the agent's ability to reason over its knowledge and behavior correctly, and the criticality of the task at hand. Autonomy levels and human interactions should therefore not be applied at the same level for all tasks, in all contexts, and across all phases of the mission, but instead need to be customized according to actions, context, phase, and even human preferences. This introduces the need for a systematic requirements elicitation process to explore the knowledge needs of humans and agents, and identify points at which humans can interact with the agents' autonomous behavior.

To support the elicitation, analysis, and specification of human multi-agent interactions, we developed a set of probing questions [29], [30]. These questions can be used to elicit requirements for each human multi-agent interaction point from system stakeholders. Probing questions are not necessarily easy to answer especially as human multi-agent interactions represent an emergent area of study with unknown unknowns [31]. Answering the questions therefore requires a rigorous and systematic requirements engineering elicitation and analysis process that includes brainstorming, interviews, immersive prototyping, and even field-studies in order to fully discover the requirements [32], [33].

We structure our probing questions around the four types of human multi-agent interactions defined in Figure 3. These include (1) information sharing, (2) direct feedback and commands, (3) raising or lowering of autonomy levels, and (4) providing behavior rationales and explanations. We map each question to the entities of the meta-model, and then use the answers to specify the requirements for each interaction point as a conceptual model. In each case, the first question is designed to identify specific interaction points, while all subsequent questions are used to explore the details of each interaction.

A. Sharing Information

At the most basic level, humans and agents must share information with each other in order to create a common understanding of the state of the mission and its environment. We therefore start by posing two key questions concerning the exchange of information.

- PQ1: **What information** do agents or humans need to know about the state of the mission and the environment in which they operate individually or collaboratively? [Knowledge, Role, AutonomousDecisions]
- PQ2: When and how will these agents or humans share or acquire information? [Knowledge, Information, Role]

By default, the system must be designed such that information is shared freely across humans and agents. For example, agents acquire knowledge about the environment and the state of the mission through their sensors (e.g., victim detected or wind velocity 20 mph) and through decisions they make (e.g., UAV-1 is tracking a detected victim). They share this information with other active agents and with humans on the ground. However, above and beyond this general exchange of information, we must explore additional explicit interaction points between humans and agents in order to understand the system's requirements.

B. Feedback and Commands

All five of the scenarios in Section II-B introduce the possibility of a human offering feedback or even direct commands. To elicit a more complete list of interaction points, we ask the following question:

- PQ3: When should a **human intervene** by providing direct feedback or commands to multiple agents? [IssuedCommand, AutonomousDecision, ProvidedInformation]

We then ask additional probing questions to explore each of the identified intervention points:

- PQ4: What **triggers** the feedback or command? (e.g., solicited by UAV, triggered by a specific event, or offered by the human operator based on his/her general awareness) [AutonomousDecision, Information]
- PQ5: What **information** should be provided in the feedback or command? (e.g., knowledge of the scene, permission to perform a specific task, a hint) [Information]
- PQ6: How should the agent **respond to the feedback**? (e.g., update its situational awareness, obey the command regardless of its current environmental knowledge) [Role, AutonomousDecisions]

C. Providing behavioral rationales

Scenarios S4 and S5 provided clear examples in which a UAV needed to explain its behavior. To identify other such interaction points we pose the following question:

- PQ7: In what concrete situations would humans require agents to explain themselves? [AutonomousDecision]

The following questions are then posed for each situation in which the agent is expected to explain its behavior.

- PQ8: Why does the agent need to **explain itself** at this collaboration point? (e.g., unexpected behavior) [Role]
- PQ9: What **information** needs to be included in the explanation? (e.g., current task, goals, actions, rationales) [Information]
- PQ10: Under what circumstances might the human choose to override the agent's decision based on its explanation? If so, what would those overrides look like? (e.g., feedback/command, or lowering of autonomy levels.) [HumanInteraction, ProvidedInformation, IssuedCommand, ChangedPermission]

D. Raising or Lowering of Autonomy Levels

Scenarios S4 and S5 also provide examples where a human operator may wish to raise or lower autonomy levels. To identify such intervention points we pose the following question:

- PQ11: When and where do the agents exhibit **autonomous decision-making behavior**? [Role, AutonomousDecision]

Each identified intervention point is then explored through the following questions:

- PQ12: **What information** do the agents need in order to exhibit the autonomous behavior? [Information]
- PQ13: Under **normal operating conditions**, what decisions should the agent be able to make autonomously? [AutonomousDecision]
- PQ14: What **constraints** on the agent's autonomy are introduced by issues related to safety, ethics, regulatory requirements, or human trust? (e.g., FAA Part 107 regulations prohibit night-time flight without an explicit waiver) [Permission]
- PQ15: How is the **autonomy suppressed or increased** at this interaction point? (e.g., modifying the confidence threshold for automatically tracking a potential victim, disabling/enabling the ability to track without permission, disabling/enabling the ability for a UAV to determine its ideal altitude and velocity during a search – or altering the range of allowed values.) [Role, ChangedPermission]
- PQ16: Are there circumstances in which the human needs to make run-time decisions about suppressing or raising autonomy (i.e., human interaction is required) vs. clearly defined rules by which the autonomy levels can be automatically raised and lowered? [Permission, ChangedPermission]
- PQ17: When autonomy is suppressed or increased what **extra support structures** would be needed, if any, for the emergency responders? (e.g., the operator manually pilots multiple UAVs and additional 360° views are needed). [Role]

E. Constructing Requirements Models

For each identified human multi-agent interaction point we specify requirements for the interaction by constructing a

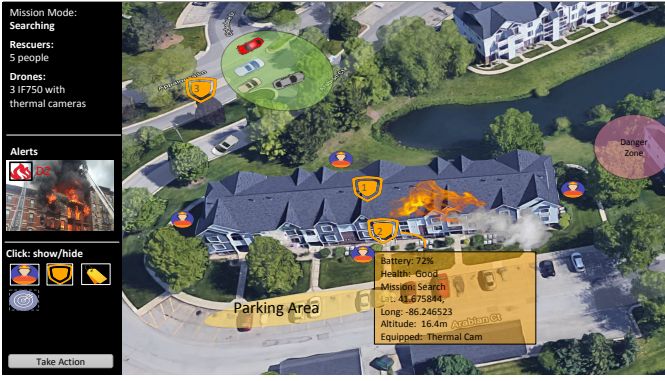


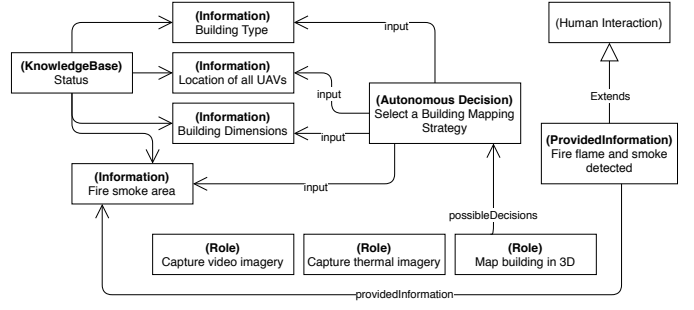
Fig. 7: A visionary prototype that we used during brainstorming meetings with Fire Fighters to trigger ideas about the use of DroneResponse in fighting structural fires.

example, depending upon the size and layout of the building, weather conditions, and the number of available UAVs, they could work independently on different parts (sides, roof) of the building, they could prioritize specific areas, fly around the building in either direction, or even work together on a single section at distinct altitudes. In the scenario that we model, the UAVs devise a specific mapping plan; however, firefighters observe smoke coming from a different area of the building, update the knowledge base, and this leads to the UAVs redesigning their strategy. In this example, the firefighters do not issue a direct command, but instead provide additional information and allow the UAVs to autonomously adapt their plans. In this example, one of the UAVs assumes a new role of using thermal imagery to search for victims through windows in the area at which smoke has been detected.

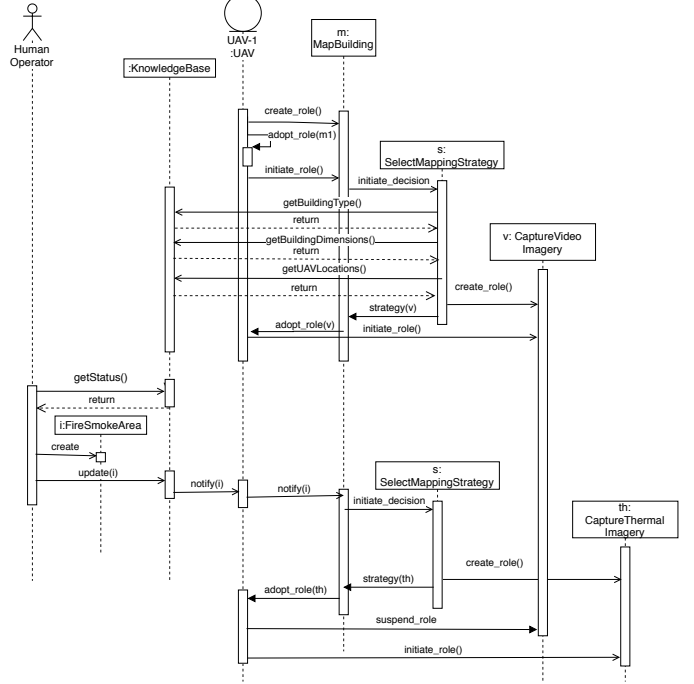
The probing questions enable us to explore this type of scenario. *PQ11*, *PQ12*, and *PQ13* identify the required *AutonomousDecisions* and the required *Information* to create the 3D model of the building autonomously. *PQ3* and *PQ4* elicit human multi-UAV interaction points such as *fire smoke detection by humans* while UAVs are engaged in mapping the building. *PQ6* identifies potential flight adaptation patterns and roles assumed by the UAVs after receiving updated information about the smoke. Answers to the probing questions lead us to construct the conceptual model and sequence diagram depicted in Figure 8.

VII. THREATS TO VALIDITY

There are several threats to validity for our approach. First, we have applied the probing questions retrospectively to construct the M1 models described in the fire-surveillance example; however, we answered the questions based on information gathered through a series of brainstorm meetings with firefighters. In the next phase of our work, we will further evaluate the questions in live requirements elicitation sessions. Second, we developed our meta-model based on five use-cases primarily developed by our own research group in collaboration with our local fire fighters. We then demonstrated its generalizability using an additional use case that



a: Conceptual model showing entities involved in the human multi-agent interaction.



b: Sequence diagram showing human-agent collaborations.

Fig. 8: Multiple UAVs collaborate to create the 3D model of the building. When the human operator shares information about smoke observed at one end of the building, UAV-1 starts capturing thermal images to search for victims through the smoke.

we developed. Our approach needs to be evaluated on use-cases elicited from diverse groups of emergency responders. Finally, our approach currently ends at the modeling stage. To fully evaluate the usefulness of the model and the probing questions, we need to implement and integrate the modeled interactions within our deployed system. We are currently working towards developing the required infrastructure such as AI vision models, on-board analysis and reasoning framework to support autonomous capabilities of the UAVs, and will then evaluate the extent to which our approach produces a viable design for use with physical UAVs.

VIII. RELATED WORK

The effectiveness of the HotL is highly dependent upon the human multi-agent interaction mechanisms built into the system as well as the flexibility of the autonomy models. To this end, several researchers have explored techniques for exposing the intent, actions, plans and associated rationales of an autonomous agent [34], while other researchers have explored ways to improve overall performance by dynamically adapting agents' autonomy levels based on the estimated cognitive workload of the human participants; however, they also observed that frequent changes in autonomy levels reduced situational awareness and forced operators to continually reevaluate the agents' behavior [35].

Furthermore, systems that use AI techniques to support autonomy often lack adequate explanations of the autonomous behavior which can negatively impact achievement of mission goals [36] and reduce trust in the system. Therefore, several of our PQs are specifically designed to explore the explainable aspects of a HotL system. Guizzardi argues that RE techniques can be applied in the design of AI systems, such as driverless cars and autonomous weapons, to ensure that they comply to ethical principles and codes [37]. Gamification is a popular technique for gathering and validating the requirements of a cyber-physical system [38]. Wiesner et al. [39] engaged stakeholders in a simulated game under different operational conditions to discover the limitations of the existing requirements and to support the ideation of possible new services. Fischer also uses a multi-player mixed-reality game to generate requirements for interaction and coordination within rich and 'messy' real-world socio-technical settings [40]. However, Hyrynsalmi discusses limitations of gamification techniques [41], for example, users focusing on winning the 'game' instead of the challenges of interacting with the system [42]. The gamification approach also requires a significant upfront development effort and proves insufficient to explore the unknown unknowns of the system. Our work takes a more formal approach to elicit requirements using a concrete meta-model and PQs that focus on the human interaction aspects of the multi-agent HotL systems.

IX. CONCLUSION

This paper describes the model-driven analysis and specification of human multi-agent interaction requirements for a human-on-the-loop system. The human multi-agent interaction types, the proposed meta-model, and the structured probing questions assist in modeling and formally specifying the complex human multi-agent interactions. We have demonstrated its use through formally specifying human interaction and intervention points for two distinct scenarios in which multiple semi-autonomous UAVs are deployed in emergency response missions. Our future work will involve implementing and evaluating our models with first-responders with physical UAVs in outdoor field-tests.

X. ACKNOWLEDGEMENT

The work described in this paper is partially funded by the USA National Science Foundation grant CNS-1931962 .

REFERENCES

- [1] J. Torresen, "A review of future and ethical perspectives of robotics and ai," *Frontiers in Robotics and AI*, vol. 4, p. 75, 2018.
- [2] M. Carpentiero, L. Gugliemetti, M. Sabatini, and G. B. Palmerini, "A swarm of wheeled and aerial robots for environmental monitoring," in *2017 IEEE 14th International Conference on Networking, Sensing and Control (ICNSC)*. IEEE, 2017, pp. 90–95.
- [3] S.-J. Chung, A. A. Paranjape, P. Dames, S. Shen, and V. Kumar, "A survey on aerial swarm robotics," *IEEE Transactions on Robotics*, vol. 34, no. 4, pp. 837–855, 2018.
- [4] Z. Hu, K. Wan, X. Gao, Y. Zhai, and Q. Wang, "Deep reinforcement learning approach with multiple experience pools for uav's autonomous motion planning in complex unknown environments," *Sensors*, vol. 20, no. 7, p. 1890, 2020.
- [5] M. Kläs and A. M. Vollmer, "Uncertainty in machine learning applications: A practice-driven classification of uncertainty," in *International Conference on Computer Safety, Reliability, and Security*. Springer, 2018, pp. 431–438.
- [6] S. Nahavandi, "Trusted autonomy between humans and robots: Toward human-on-the-loop in robotics and autonomous systems," *IEEE Systems, Man, and Cybernetics Magazine*, vol. 3, no. 1, pp. 10–17, 2017.
- [7] J. E. Fischer, C. Greenhalgh, W. Jiang, S. D. Ramchurn, F. Wu, and T. Rodden, "In-the-loop or on-the-loop? interactional arrangements to support team coordination with a planning agent," *Concurrency and Computation: Practice and Experience*, p. e4082, 2017.
- [8] R. K. Bellamy, S. Andrist, T. Bickmore, E. F. Churchill, and T. Erickson, "Human-agent collaboration: Can an agent be a partner?" in *Proceedings of the 2017 CHI Conference Extended Abstracts on Human Factors in Computing Systems*, 2017, pp. 1289–1294.
- [9] J. Cleland-Huang and A. Agrawal, "Human-drone interactions with semi-autonomous cohorts of collaborating drones," in *Interdisciplinary Workshop on Human-Drone Interaction (iHDI 2020)*, *CHI '20 Extended Abstracts*, 26 April 2020, Honolulu, HI, US, 2020.
- [10] P. A. Hancock and S. F. Scallen, "Allocating functions in human-machine systems," 1998.
- [11] Z. Xu, J. Yang, C. Peng, Y. Wu, X. Jiang, R. Li, Y. Zheng, Y. Gao, S. Liu, and B. Tian, "Development of an uas for post-earthquake disaster surveying and its application in ms7.0 lushan earthquake, sichuan, china," *Computers & Geosciences*, vol. 68, pp. 22 – 30, 2014. [Online]. Available: <http://www.sciencedirect.com/science/article/pii/S0098300414000788>
- [12] E. D. Beni, M. Cantarero, and A. Messina, "Uavs for volcano monitoring: A new approach applied on an active lava flow on mt. etna (italy), during the 27 february–02 march 2017 eruption," *Journal of Volcanology and Geothermal Research*, vol. 369, pp. 250 – 262, 2019. [Online]. Available: <http://www.sciencedirect.com/science/article/pii/S0377027318301756>
- [13] G. Dooly, E. Omerdic, J. Coleman, L. Miller, A. Kaknjo, J. Hayes, J. Braga, F. Ferreira, H. Conlon, H. Barry, J. Marcos-Olaya, T. Tuohy, J. Sousa, and D. Toal, "Unmanned vehicles for maritime spill response case study: Exercise cathach," *Marine Pollution Bulletin*, vol. 110, no. 1, pp. 528 – 538, 2016. [Online]. Available: <http://www.sciencedirect.com/science/article/pii/S0025326X16301242>
- [14] M. Fleck, "Usability of lightweight defibrillators for uav delivery," in *Proceedings of the 2016 CHI Conference Extended Abstracts on Human Factors in Computing Systems*, ser. CHI EA '16. New York, NY, USA: Association for Computing Machinery, 2016, p. 3056–3061. [Online]. Available: <https://doi.org/10.1145/2851581.2892288>
- [15] N. Athanasis, M. Themistocleous, K. Kalabokidis, and C. Chatzitheodorou, "Big data analysis in uav surveillance for wildfire prevention and management," in *Information Systems*, M. Themistocleous and P. Rupino da Cunha, Eds. Cham: Springer International Publishing, 2019, pp. 47–58.
- [16] D. Tezza and M. Andujar, "The state-of-the-art of human-drone interaction: A survey," *IEEE Access*, vol. 7, pp. 167 438–167 454, 2019.
- [17] M. Funk, "Human-drone interaction: Let's get ready for flying user interfaces!" *Interactions*, vol. 25, no. 3, pp. 78–81, 2018. [Online]. Available: <https://doi.org/10.1145/3194317>

- [18] J. R. Cauchard, J. L. E. K. Y. Zhai, and J. A. Landay, "Drone & me: An exploration into natural human-drone interaction," in *Proceedings of the 2015 ACM International Joint Conference on Pervasive and Ubiquitous Computing*, ser. UbiComp '15. New York, NY, USA: Association for Computing Machinery, 2015, p. 361–365. [Online]. Available: <https://doi.org/10.1145/2750858.2805823>
- [19] A. Agrawal, S. J. Abraham, B. Burger, C. Christine, L. Fraser, J. M. Hoeksema, S. Hwang, E. Travník, S. Kumar, W. Scheirer, J. Cleland-Huang, M. Vierhauser, R. Bauer, and S. Cox, "The next generation of human-drone partnerships: Co-designing an emergency response system," in *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*, ser. CHI '20. New York, NY, USA: Association for Computing Machinery, 2020, p. 1–13. [Online]. Available: <https://doi.org/10.1145/3313831.3376825>
- [20] J. Cleland-Huang, M. Vierhauser, and S. Bayley, "Dronology: an incubator for cyber-physical systems research," in *Proc. of the 40th Int'l Cong. on Software Engineering: New Ideas and Emerging Results*, 2018, pp. 109–112. [Online]. Available: <https://doi.org/10.1145/3183399.3183408>
- [21] M. Vierhauser, J. Cleland-Huang, S. Bayley, T. Krismayer, R. Rabiser, and P. Grünbacher, "Monitoring CPS at runtime - A case study in the UAV domain," in *Proc. of the 44th Euromicro Conf. on Software Engineering and Advanced Applications, SEAA 2018, Prague, Czech Republic, August 29-31, 2018*, 2018, pp. 73–80. [Online]. Available: <https://doi.org/10.1109/SEAA.2018.00022>
- [22] M. Wooldridge, "Agent-based software engineering," *IEEE Proceedings-software*, vol. 144, no. 1, pp. 26–37, 1997.
- [23] A. Pokahr, L. Braubach, and W. Lamersdorf, "Jadex: A bdi reasoning engine," in *Multi-agent programming*. Springer, 2005, pp. 149–174.
- [24] J. Cleland-Huang and M. Vierhauser, "Discovering, analyzing, and managing safety stories in agile projects," in *26th IEEE International Requirements Engineering Conference*, 2018, pp. 262–273. [Online]. Available: <https://doi.org/10.1109/RE.2018.00034>
- [25] P. Scharre and M. Horowitz, *An introduction to autonomy in weapon systems*. Center for a New American Security, 2015.
- [26] M. R. Endsley, "Toward a theory of situation awareness in dynamic systems," *Human Factors*, vol. 37, no. 1, pp. 32–64, 1995. [Online]. Available: <https://doi.org/10.1518/001872095779049543>
- [27] R. Loftin, B. Peng, J. MacGlashan, M. L. Littman, M. E. Taylor, J. Huang, and D. L. Roberts, "Learning behaviors via human-delivered discrete feedback: modeling implicit feedback strategies to speed up learning," *Autonomous agents and multi-agent systems*, vol. 30, no. 1, pp. 30–59, 2016.
- [28] A. L. Thomaz, C. Breazeal *et al.*, "Reinforcement learning with human teachers: Evidence of feedback and guidance with implications for learning performance," in *Aaai*, vol. 6. Boston, MA, 2006, pp. 1000–1005.
- [29] P. R. Anish, B. Balasubramaniam, A. Sainani, J. Cleland-Huang, M. Daneva, R. J. Wieringa, and S. Ghaisas, "Probing for requirements knowledge to stimulate architectural thinking," in *Proceedings of the 38th International Conference on Software Engineering*, 2016, pp. 843–854.
- [30] R. E. Miller, *The Quest for Software Requirements*. MavenMark Books, USA, 2009.
- [31] A. Sutcliffe and P. Sawyer, "Requirements elicitation: Towards the unknown unknowns," in *21st IEEE International Requirements Engineering Conference, RE 2013, Rio de Janeiro-RJ, Brazil, July 15-19, 2013*. IEEE Computer Society, 2013, pp. 92–104. [Online]. Available: <https://doi.org/10.1109/RE.2013.6636709>
- [32] S. Robertson, "Requirements trawling: techniques for discovering requirements," *Int. J. Hum. Comput. Stud.*, vol. 55, no. 4, pp. 405–421, 2001. [Online]. Available: <https://doi.org/10.1006/ijhc.2001.0481>
- [33] A. G. Sutcliffe, "Requirements engineering for complex collaborative systems," in *5th IEEE International Symposium on Requirements Engineering (RE 2001), 27-31 August 2001, Toronto, Canada*. IEEE Computer Society, 2001, pp. 110–119. [Online]. Available: <https://doi.org/10.1109/ISRE.2001.948550>
- [34] J. Y. Chen, S. G. Lakhmani, K. Stowers, A. R. Selkowitz, J. L. Wright, and M. Barnes, "Situation awareness-based agent transparency and human-autonomy teaming effectiveness," *Theoretical issues in ergonomics science*, vol. 19, no. 3, pp. 259–282, 2018.
- [35] J. Heard, J. Fortune, and J. A. Adams, "Sahrta: A supervisory-based adaptive human-robot teaming architecture," *arXiv preprint arXiv:2003.05823*, 2020.
- [36] I. Stoica, D. Song, R. A. Popa, D. Patterson, M. W. Mahoney, R. Katz, A. D. Joseph, M. Jordan, J. M. Hellerstein, J. E. Gonzalez *et al.*, "A berkeley view of systems challenges for ai," *arXiv preprint arXiv:1712.05855*, 2017.
- [37] R. Guizzardi, G. Amaral, G. Guizzardi, and J. Mylopoulos, "Ethical requirements for ai systems," in *Canadian Conference on Artificial Intelligence*. Springer, 2020, pp. 251–256.
- [38] P. Lombriser, F. Dalpiaz, G. Lucassen, and S. Brinkkemper, "Gamified requirements engineering: model and experimentation," in *International Working conference on requirements engineering: foundation for software quality*. Springer, 2016, pp. 171–187.
- [39] S. Wiesner, J. B. Hauge, F. Haase, and K.-D. Thoben, "Supporting the requirements elicitation process for cyber-physical product-service systems through a gamified approach," in *IFIP International Conference on Advances in Production Management Systems*. Springer, 2016, pp. 687–694.
- [40] J. E. Fischer, W. Jiang, A. Kerne, C. Greenhalgh, S. D. Ramchurn, S. Reece, N. Pantidi, and T. Rodden, "Supporting team coordination on the ground: requirements from a mixed reality game," in *COOP 2014-Proceedings of the 11th International Conference on the Design of Cooperative Systems, 27-30 May 2014, Nice (France)*. Springer, 2014, pp. 49–67.
- [41] S. Hyrnsalmi, J. Smed, and K. Kimppa, "The dark side of gamification: How we should stop worrying and study also the negative impacts of bringing game design elements to everywhere," in *GamiFIN*, 2017, pp. 96–104.
- [42] K. Knaving and S. Björk, "Designing for fun and play: exploring possibilities in design for gamification," in *Proceedings of the first International conference on gameful design, research, and applications*, 2013, pp. 131–134.