

Multi-UAV Adaptive Path Planning Using Deep Reinforcement Learning

Jonas Westheider, Julius Rückin, Marija Popović

Abstract—Efficient aerial data collection is important in many remote sensing applications. In large-scale monitoring scenarios, deploying a team of unmanned aerial vehicles (UAVs) offers improved spatial coverage and robustness against individual failures. However, a key challenge is cooperative path planning for the UAVs to efficiently achieve a joint mission goal. We propose a novel multi-agent informative path planning approach based on deep reinforcement learning for adaptive terrain monitoring scenarios using UAV teams. We introduce new network feature representations to effectively learn path planning in a 3D workspace. By leveraging a counterfactual baseline, our approach explicitly addresses credit assignment to learn cooperative behaviour. Our experimental evaluation shows improved planning performance, i.e. maps regions of interest more quickly, with respect to non-counterfactual variants. Results on synthetic and real-world data show that our approach has superior performance compared to state-of-the-art non-learning-based methods, while being transferable to varying team sizes and communication constraints.

I. INTRODUCTION

Efficient aerial data collection is crucial for mapping and monitoring phenomena on the Earth's surface. Unmanned aerial vehicles (UAVs) provide a flexible, labour-, and cost-efficient solution for remote sensing applications such as precision agriculture [1–3], wildlife conservation [4], and search and rescue [5, 6]. For monitoring large terrains, replacing a single UAV with a multi-UAV system can improve spatial coverage, versatility, and robustness to individual failures at lower overall cost [7]. However, to fully unlock its potential, a key challenge is planning UAV paths cooperatively in complex environments, given on-board constraints on runtime efficiency and communication.

This paper addresses active data collection using a team of UAVs in terrain monitoring scenarios. Our goal is to map an initially unknown, non-homogeneous binary target variable of interest on a 2D terrain, e.g. crop infestations in an agricultural scenario or to-be-rescued victims in a disaster scenario, using image measurements taken by the UAVs. We tackle the problem of multi-agent *informative path planning* (IPP): we plan information-rich paths for the UAVs to cooperatively gather sensor data subject to constraints on energy, time, or distance. Our motivation is to allow the UAVs to *adaptively* monitor the terrain in areas of interest where information value is high.

Traditional approaches for data collection are *non-adaptive* and rely on static, predefined paths. In multi-UAV coverage

This work was funded by the Deutsche Forschungsgemeinschaft (DFG, German Research Foundation) under Germany's Excellence Strategy - EXC 2070 - 390732324. Authors are with the Cluster of Excellence PhenoRob, Institute of Geodesy and Geoinformation, University of Bonn. Corresponding: jwesthei@uni-bonn.de.

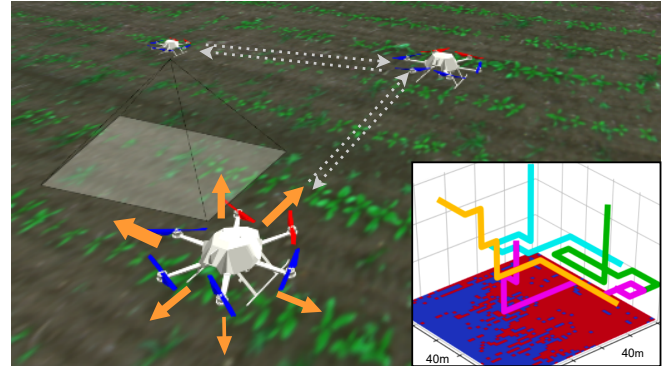


Fig. 1: Our RL-based approach applied in a multi-UAV surface temperature mapping scenario. The UAVs take images of the terrain (white transparent) and communicate them (grey dashed arrows). Based on locally available information, each UAV decides from a set of actions (orange arrows) where to take the next measurement. Inset image: resulting trajectories for 4 deployed UAVs (different colours). By planning paths cooperatively, our approach enables adaptively mapping warm (red) areas of interest on the field.

path planning [8], the terrain is equally partitioned and a sweep pattern is assigned to each UAV. The main drawback of such methods is that they assume uniformly distributed target variables of interest, e.g. anomalies, hotspots, victims, and do not allow for targeted inspection of specific areas. To address this, IPP methods [1, 9–13] have been proposed enabling online decision-making based on incoming information. However, the runtime of these strategies usually scales exponentially with the planning horizon, since they rely on evaluating many candidate paths online. In team scenarios, reasoning about the other UAVs' behaviours causes the number of evaluations to further grow exponentially with the team size, which leads to intractable runtime complexity.

Reinforcement learning (RL) has emerged as a popular approach for efficiently learning online decision-making in robotics [11–16]. Recent works apply RL for single-agent IPP to enhance path quality and computation time for adaptive data collection [11, 14, 17, 18]. However, learning informative paths for multiple cooperating agents remains an open challenge. First studies show promising results [12, 13], but are limited to 2D action spaces and do not address the credit assignment problem [19], i.e. how much each agent contributes to the overall team performance, which adversely impacts cooperation capabilities.

The main contribution of this paper is a novel multi-agent deep RL-based IPP approach for adaptive terrain monitoring scenarios using UAV teams. Bridging the gap between recent advances in RL and robotic applications, we build upon counterfactual multi-agent policy gradients

(COMA) [20] to explicitly address the credit assignment problem in cooperative IPP. Our approach supports decentralised on-board decision-making and achieves cooperative 3D path planning with variable team sizes. We show that (i) our designed network input representations are effective for multi-agent IPP in 3D action spaces; (ii) our COMA-based algorithm accounts for credit assignment resulting in improved planning performance; and (iii) our RL-based approach improves planning performance compared to non-learning-based methods across varying team sizes and communication constraints without re-training. To back up these claims, we demonstrate the performance of our approach using synthetic and real-world data in a thermal hotspot mapping scenario. We will open-source our code at: <https://github.com/dmar-bonn/ipp-marl>.

II. RELATED WORK

Our work brings together multi-robot IPP for autonomous data collection and advances in multi-agent RL. This section overviews relevant literature for both, distinguishing between single- and multi-robot settings.

Informative path planning has been widely studied for efficient data gathering using autonomous robots [1, 2, 9, 11–13, 17, 21]. *Non-adaptive* coverage planners [8] monitor a terrain exhaustively based on pre-defined paths. In contrast, we focus on *adaptive* strategies replanning paths online based on incoming sensor data. Hollinger et al. [9] study the benefit of adaptive online replanning to maximise the information value of sensor measurements given a limited mission budget for single-vehicle underwater inspection tasks. Recent works [21, 22] propose optimisation-based IPP methods for single-UAV monitoring. However, these methods require computationally costly evaluations of candidate paths' expected information value. Extending them naïvely to multiple UAVs yields exponentially scaling complexity which restricts applicability to online applications.

Deploying cooperative robotic teams is beneficial for monitoring tasks over large terrains such as post-disaster assessment [5, 6] and precision agriculture [1–3]. We focus on decentralised planning strategies, where robots make locally informed decisions on-board to enable robustness to individual failures and scalability with the team size. Best et al. [23] introduce a decentralised variant of Monte Carlo tree search over a joint probability distribution of action sequences to plan individual paths for active perception. Similar to our multi-UAV setup, Albani et al. [1] and Carbone et al. [2] propose methods for monitoring areas of interest in crop fields. The former use biology-inspired swarm behavior for adaptive IPP. Despite promising results, their method relies on manual problem-specific parameter tuning.

An alternative line of work leverages learning-based methods for active data collection. Following the paradigm of centralised training and decentralised policy execution, Li et al. [24] learn inter-robot communication policies to exchange local robot information. Tzes et al. [10] propose a more general multi-robot framework relying on graph neural networks. Both approaches are trained via supervised

imitation learning and require an expert to learn from. We learn the desired team behaviour using deep RL, which enables flexible, foresighted planning in various applications.

Reinforcement learning (RL) is increasingly utilised in robotics and UAV applications [11, 14, 17, 18], combining remarkable planning results with little computational effort. Recently, RL has also been applied to IPP to efficiently replan robotic paths online. Chen et al. [14] develop a graph-based deep RL method for exploration, selecting map frontiers that reduce map uncertainty and travel time. However, their approach is limited to 2D workspaces, while we consider 3D planning. Other works reward agent actions that lead to high information gain [18] and map uncertainty reduction in target regions [17]. In a similar problem setup to ours, Rückin et al. [11] propose an IPP method combining deep RL and sampling-based planning for adaptive UAV terrain monitoring in 3D workspaces. Naïvely extending single-agent algorithms to multi-agent setups incurs exponentially growing complexity with the number of agents.

Multi-agent RL for IPP is a relatively unexplored research area. Existing approaches for cooperative team applications do not account for the credit assignment problem [19], i.e. do not discriminate the contribution of one agent to the overall team performance. Most works address related applications including navigation [15], target assignment [25], and coverage planning [26]. Similar to us, Bayerlein et al. [13] propose a multi-agent RL-based IPP approach that maximises harvested data without inter-agent communication. Viseras and Garcia [12] allow agents to exchange information via a communication module. Both works are limited to constant UAV altitudes ignoring potentially varying sensor noise with altitude. Recently, Luis et al. [16] proposed a deep Q-learning algorithm that supports learning cooperation by penalising redundant measurements based on the inter-UAV distance. Their reward design is tailored to pure exploration instead of adaptively monitoring areas of interest.

These works independently assign hand-engineered individual agent rewards. Thus, they do not fully account for the cooperative nature of IPP and can fail when greedily aiming for a high individual reward contradicts the desired team behaviour. In contrast, we propose a new approach solely relying on generally applicable global rewards for the whole UAV team. We adapt the counterfactual multi-agent (COMA) RL algorithm [20] to active robotic data collection in 3D workspaces. This way, we explicitly assign credit to individual agents during training. Our experimental results emphasise the need for explicit credit assignment to achieve cooperative behaviour and verify improved IPP performance.

III. PROBLEM STATEMENT

We consider a team of homogeneous UAVs monitoring a static flat terrain in an obstacle-free environment. The goal is to plan information-rich UAV paths on-the-fly as new measurements arrive to efficiently map regions of interest on the terrain given a finite mission budget. We briefly describe the general multi-agent IPP problem, our mapping strategy, and how to quantify information value for our RL approach.

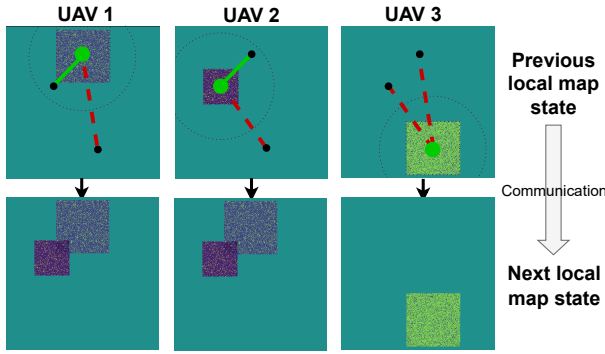


Fig. 2: Inter-UAV communication. The UAVs (filled circles) exchange their current measurements (square footprints) with each other when closer than a limited communication range (black dotted circle). Green and dotted red lines indicate in-range and out-of-range communications, respectively. The UAVs use the received measurements to update their local map states (top to bottom row).

A. Multi-Agent Informative Path Planning

We address the problem of multi-agent informative path planning (IPP) optimising an information-theoretic criterion $I : \Psi^N \rightarrow \mathbb{R}^+$ over all N UAV paths $\psi = \{\psi_1, \dots, \psi_N\}$, where Ψ is the set of all possible individual UAV paths:

$$\psi^* = \operatorname{argmax}_{\psi \in \Psi^N} I(\psi), \text{ s.t. } C(\psi_i) \leq B \quad \forall i \in \{1, \dots, N\}. \quad (1)$$

Each set of UAV paths ψ is composed of individual paths $\psi_i = (\mathbf{p}_i^0, \dots, \mathbf{p}_i^B) \in \Psi$ of length B with 3D measurement positions $\mathbf{p}_i^t \in \mathbb{R}^3$ in an equi-distant grid $\mathcal{P}$ of resolution r_P over multiple altitudes above the terrain. The function $C : \Psi \rightarrow \mathbb{R}^+$ maps a UAV path ψ_i to its associated execution cost; in our work, a maximum number $B \in \mathbb{N}^+$ of measurements taken along a path ψ_i .

B. Terrain Mapping

The UAVs map the terrain by taking images using downward-facing cameras. The image information is processed, e.g. by semantic segmentation, communicated, and fused into a terrain map. A measurement z_i^t taken by UAV i at time step t is a likelihood estimate projected onto the flat terrain. We consider binary per-pixel classification and sequentially fuse z_i^t using probabilistic occupancy grid mapping [27]. Each UAV i stores a local posterior map belief $\mathcal{M}_i$ discretised into M grid cells $\mathcal{M}_i^j$ of resolution r_M . The altitude of the current measurement position $\mathbf{p}_i^t$ determines the mapped field of view and the mapping resolution. We align r_M with the mapping resolution at the lowest altitude and upsample higher-altitude measurements to r_M to fuse heterogeneous mapping resolutions together in $\mathcal{M}_i$. Similar to Popović et al. [21], we leverage an altitude-dependent sensor model specifying $p(z_i^t | \mathcal{M}_i^j, \mathbf{p}_i^t)$ to account for non-uniform sensor measurement noise and update a UAV's local posterior map belief over each $\mathcal{M}_i^j$ at time step t .

Fig. 2 visualises the inter-UAV communication protocol. All UAVs i send their measurements z_i^t to another UAV k and receive measurements z_k^t to update the individual local map beliefs over $\mathcal{M}_i$. UAVs i and k share measurements, if

$\|\mathbf{p}_i^t - \mathbf{p}_k^t\|_2 \leq D$, where $D \in \mathbb{R}^+$ is a radius approximating a range-limited communication channel. The communication message contains the UAV identifier i , its position $\mathbf{p}_i^t$, and its collected measurement z_i^t . The receiving UAV k utilises z_i^t and $\mathbf{p}_i^t$ to update its local map belief over $\mathcal{M}^k$.

C. Utility Definition for Adaptive Terrain Monitoring

Our goal is to plan future measurement positions $\psi_i^{t+1} = (\mathbf{p}_i^{t+1}, \dots, \mathbf{p}_i^B)$ for each UAV i at time step t from which the next measurements $\{z_i^{t+1}, \dots, z_i^B\}$ maximally reduce the uncertainty of the current posterior belief over a global map $\mathcal{M}$. The global map $\mathcal{M}$ contains measurements $z_i^{0:t}$ taken by the N UAVs up to the current time step t . We quantify the uncertainty reduction of a set of measurements $z^{t+1} = \{z_1^{t+1}, \dots, z_N^{t+1}\}$ taken at the next time step $t+1$ by computing the entropy reduction over $\mathcal{M}$ after fusing z^{t+1} . Then, the information criterion $I(\psi^{t+1})$ is computed as the summed entropy reduction along paths ψ^{t+1} :

$$I(\psi^{t+1}) = \sum_{m=0}^{B-t-1} H(\mathcal{M} | z^{0:t+m}, \mathbf{p}^{0:t+m}) - H(\mathcal{M} | z^{0:t+m+1}, \mathbf{p}^{0:t+m+1}). \quad (2)$$

We consider adaptive mapping with one interesting target class, i.e. one class that holds information value, and a complement class. Thus, the map entropy $H(\mathcal{M} | z^{0:t}, \mathbf{p}^{0:t})$ at time step t is weighted by the importances of the interesting and the uninteresting classes:

$$\begin{aligned} H(\mathcal{M} | z^{0:t}, \mathbf{p}^{0:t}) &= \sum_{j=1}^M H(\mathcal{M}^j | z^{0:t}, \mathbf{p}^{0:t}) \\ &= - \sum_{j=1}^M W(\mathcal{M}^j) p(\mathcal{M}^j | z^{0:t}, \mathbf{p}^{0:t}) \log(p(\mathcal{M}^j | z^{0:t}, \mathbf{p}^{0:t})) \\ &\quad + W(\overline{\mathcal{M}^j}) p(\overline{\mathcal{M}^j} | z^{0:t}, \mathbf{p}^{0:t}) \log(p(\overline{\mathcal{M}^j} | z^{0:t}, \mathbf{p}^{0:t})), \end{aligned} \quad (3)$$

where $p(\overline{\mathcal{M}^j} | z^{0:t}, \mathbf{p}^{0:t}) = 1 - p(\mathcal{M}^j | z^{0:t}, \mathbf{p}^{0:t})$. The importance weighting $W(\mathcal{M}^j)$ is defined as:

$$W(\mathcal{M}^j) = \begin{cases} w_1 & \text{if } p(\mathcal{M}^j | z^{0:t}, \mathbf{p}^{0:t}) > 0.5, \\ w_2 & \text{if } p(\mathcal{M}^j | z^{0:t}, \mathbf{p}^{0:t}) < 0.5, \\ 0.5 & \text{else,} \end{cases} \quad (4)$$

where $w_1, w_2 \geq 0$ and $w_1 + w_2 = 1$. The importance weights w_1, w_2 encourage the UAVs to plan paths ψ^{t+1} targeting potentially interesting regions, i.e. regions believed to contain the target class according to the current posterior map belief $p(\mathcal{M}^j | z^{0:t}, \mathbf{p}^{0:t})$.

IV. OUR APPROACH

We present our novel multi-agent RL-based IPP approach for UAV teams. Our goal is to plan UAV paths in a 3D workspace to achieve cooperative adaptive terrain mapping. As shown in Fig. 3, we train agents offline based on global terrain information using RL to learn UAV paths in a centralised way. A key aspect is the integration of a counterfactual baseline, allowing us to estimate each agent's mapping contribution to the overall team performance and

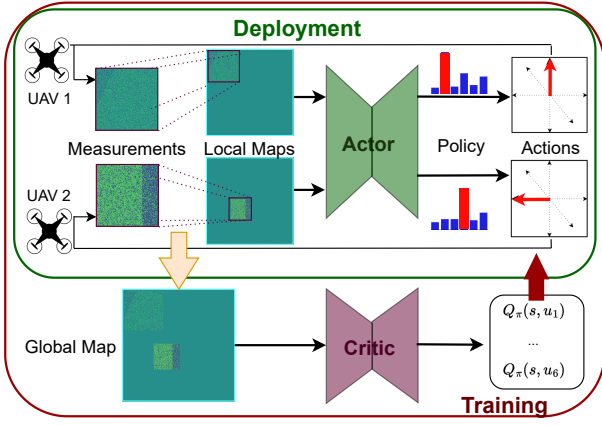


Fig. 3: Overview of our approach. At each time step during a mission, each UAV takes a measurement and updates its local map state. The local map is input to an actor network, which outputs a policy from which an action is sampled. During training, a centralised critic network is additionally trained using global map information and outputs Q-values for each action from the current state, i.e. the expected future return.

improve cooperation. During a mission, we leverage the trained agent behaviour and deploy a fully decentralised system to replan informative measurement positions online.

A. RL for Sequential Decision-Making

We formulate the multi-UAV IPP task as a sequential decision-making problem for a team of agents and train it using RL. The UAVs execute missions, where at each time step t each UAV i simultaneously plans its next measurement position $\mathbf{p}_i^{t+1}$ based on the current local on-board state ω_i^t . The local state ω_i^t includes the local map belief $\mathcal{M}_i$ that sequentially fuses past (communicated) measurements $z^{0:t}$ encoding the measurement history. This way, we account for partial observability induced by communication constraints and noisy sensor measurements. At each time step, all UAV actions, i.e. next measurement positions, define the joint action $\mathbf{u}^t = (u_1^t, \dots, u_N^t) \in U^N$. The UAVs receive one global team reward $R^t: S \times U^N \times S \rightarrow \mathbb{R}$ quantifying the joint information value of mapped measurements $z^{t+1} = \{z_1^t, \dots, z_N^t\}$. The discounted return $G^t = \sum_{k=0}^{B-t} \gamma^k R^{t+k}$, where $\gamma \in [0, 1)$ is a discount factor, measures the information value of the team's paths from t until the mission budget is spent. Although each UAV chooses actions in a decentralised way based on ω_i^t , we evaluate the performance of the team as a whole to enforce cooperative IPP behaviour.

State. The global environment state $s^t \in S$ is defined as $s^t = \{\mathcal{M}^t, \mathbf{p}_{1:N}^t, b\}$. $\mathcal{M}^t$ captures the current global map belief $p(\mathcal{M} | z^{0:t}, \mathbf{p}_{1:N}^{0:t})$, $\mathbf{p}_{1:N}^t$ are the current agent positions, and $b \leq B$ is the remaining mission budget. s^{t+1} is the next state after the agents have moved to the next positions $\mathbf{p}_{1:N}^{t+1}$ reached by executing actions $\mathbf{u}^t$. $\mathcal{M}^{t+1}$ is updated based on measurements z^{t+1} taken at $\mathbf{p}_{1:N}^{t+1}$.

Actions. The agents move within a discrete position grid P , bounded by the environment borders and discrete minimum and maximum altitudes. Each UAV i selects an action u_i from a discrete 3D action space U containing movements

$\{up, north, east, south, west, down\}$ with a fixed step size. We prevent actions leading to movement outside of the environment or UAVs having concurrent 2D terrain coordinates. **Reward.** We explicitly design the reward function R to reflect the information criterion defined in Eq. (2) to adaptively and quickly reduce map uncertainty in target areas:

$$R^t(s^t, \mathbf{u}^t, s^{t+1}) = \alpha \frac{H(\mathcal{M}^{t+1}) - H(\mathcal{M}^t)}{H(\mathcal{M}^t)} + \beta. \quad (5)$$

We reward the weighted map entropy reduction from the current to the next map state $\mathcal{M}^{t+1}$ after mapping new measurements z^{t+1} . We normalise the reward by the current map entropy to keep its magnitude approximately constant during the mission and apply affine scaling factors α and β to improve training stability. Note that, for $\gamma = 1$, the return G^t resembles the IPP criterion in Eq. (2) up to the normalisation and scaling factors.

B. Algorithm

Our goal is to learn a policy enabling cooperative UAV team behaviour for the adaptive monitoring task. To do this, we build upon the COMA RL algorithm of Foerster et al. [20] to address our IPP task for UAV teams. COMA is an actor-critic algorithm using a centralised critic network to evaluate each agent's behaviour, described by the policy $\pi(\cdot | \omega_i^t)$ and parameterised by an actor network, and to optimise the policy accordingly. The critic evaluates the current policy π by estimating the agent's i state-action value $Q_{\pi}(s^t, (u_1^t, \dots, u_i^t, \dots, u_N^t))$ given the other agents' actions $\mathbf{u}_{-i}^t$. The critic network is trained on-policy via TD(λ) [28] to estimate the discounted return G_t , introduced in Sec. IV-A, for taking the joint action $\mathbf{u}^t$ from the current state s^t and following the current policy π afterwards. The critic uses global information s^t during training, while the actor utilises only on-board information ω_i^t to predict the next-best measurement position decentralised during both training and deployment. We leverage the counterfactual baseline to assign credit to individual agents according to their contribution to the team performance, which fosters cooperative team behaviour. The advantage A_i^t for agent i taking action u_i^t in the team's action $\mathbf{u}^t$ is:

$$A_i^t(s^t, \mathbf{u}^t) = Q_{\pi}(s^t, \mathbf{u}^t) - \sum_{u_i'^t \in U} \pi(u_i'^t | \omega_i^t) Q_{\pi}(s^t, (\mathbf{u}_{-i}^t, u_i'^t)). \quad (6)$$

The contribution of agent i to the joint state-action value $Q_{\pi}(s^t, \mathbf{u}^t)$ of the team's action $\mathbf{u}^t = (u_1^t, \dots, u_i^t, \dots, u_N^t)$ by taking action u_i^t is estimated by marginalising over all possible individual actions $u_i'^t \in U$ while keeping the other agents' actions $\mathbf{u}_{-i}^t$ fixed. For optimising $\pi(\cdot | \omega_i^t)$, we apply the policy gradient theorem using Eq. (6) and minimise:

$$\mathcal{L} = -\log \pi(u_i^t | \omega_i^t) A_i^t(s^t, \mathbf{u}^t), \quad (7)$$

using mini-batch stochastic gradient descent. For further details, we refer to Foerster et al. [20].

C. Network Architecture & Feature Design

We propose new actor and critic representations to exploit COMA in 3D robotic applications. As shown in Fig. 4, actor and critic are represented by neural networks f_{θ^π} and f_{θ^c} , parameterised by θ^π and θ^c . The actor network is conditioned on the agent i and its local state ω_i^t . Feature inputs are the agent's identifier i , the remaining mission budget b , and the following spatial inputs: (a) a position map centred around the agent's position encoding the map boundaries and the communicated other agents' positions, where the values represent the agent altitudes; (b) the local map state $\mathcal{M}_i$; (c) the weighted entropy of the local map state $H(\mathcal{M}_i | z^{0:t}, \mathbf{p}^{0:t})$ in Eq. (3); (d) the weighted entropy of the measurement $H(z_i^t | \mathbf{p}_i^t)$; and (e) the map cells currently spanned by all agents' fields of view within the communication range ('footprint map').

Our critic network receives the same input (a)-(e) and, in addition, global environment information accessible during training. Specifically, it further receives: (f) a global position map encoding all agent positions $\mathbf{p}_{1:N}^t$; (g) the global map state $\mathcal{M}$; (h) its weighted entropy $H(\mathcal{M} | z^{0:t}, \mathbf{p}^{0:t})$; (i) the map cells currently spanned by all agents' fields of view; and, to enable learning the counterfactual baseline in Eq. (6), (j) the other agents' actions. Inputs are provided in the position grid resolution r_P . We downsample the map resolution r_M by $\frac{r_P}{r_M}$ to align both resolutions. The scalar inputs i and b are expanded to constant-valued 2D feature maps.

To handle spatial information necessary for the terrain monitoring task, the networks consist of convolutional encoders and multi-layer perceptron heads predicting the policy π and Q-values Q_π , respectively. Fig. 4-Top illustrates the architecture of both networks. The actor's logits $f_{\theta^\pi}(x_i^t)$ given a collection x_i^t of feature maps (a)-(e) above, are passed through a bounded softmax function $\pi(u_i | x_i^t) = (1 - \epsilon) \text{softmax}(f_{\theta^\pi}(x_i^t)) + \frac{\epsilon}{|U|}$ to predict the stochastic policy. The hyperparameter $\epsilon \in [0, 1]$ fosters exploration during training and is set to 0 at deployment. The critic outputs Q-values $Q_\pi(s^t, (u_1^t, \dots, u_i^t, \dots, u_N^t))$ for each action $u_i^t \in U$ of agent i after the last linear layer.

D. Network Training

We train an actor and a critic network offline and utilise the trained actor online at deployment. We simulate UAV missions and learn from rewards received after fusing measurements z into the global map $\mathcal{M}$. Before each mission, we generate a new terrain of resolution r_M . The terrains are split into interesting and uninteresting regions. The split is randomly oriented and interesting regions cover between 30% to 60% of the terrain to foster generalisation. We fix the initial UAV positions and execute a mission until budget B is spent. During training, actions are sampled from the actor's policy as described in Sec. IV-C, where ϵ is linearly decreased from 0.5 to 0.02 over the first 10,000 missions.

We alternate between generating 3,000 environment interactions on-policy and policy optimisation using Eq. (7). Both networks are optimised via stochastic mini-batch gradient descent for 5 epochs using the Adam optimiser with learning

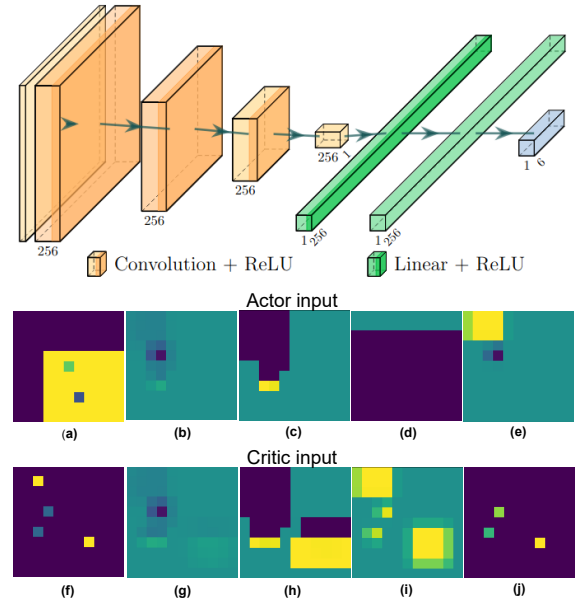


Fig. 4: Architecture and inputs for our actor and critic. Both networks consist of convolutional encoders and linear layers in the prediction head. For the actor, the output is produced by a softmax layer (grey). Both actor and critic receive local inputs ('Actor input'). The critic also receives global inputs for centralised training ('Critic input'). All inputs (a)-(j) are detailed in Sec. IV-C.

rates of $1e^{-5}$ (actor) and $1e^{-4}$ (critic), and a batch size of 600. The critic network is trained to estimate the expected return G_t applying TD(λ) with $\lambda = 0.8$ and $\gamma = 0.99$ using a target critic network copying the critic network each 30,000 environment interactions.

V. EXPERIMENTAL RESULTS

We present experiments to show the capabilities of our multi-agent RL-based IPP approach for adaptive terrain monitoring using UAV teams. Our experimental results support our claims, which are: (i) our designed network representations are effective for multi-agent IPP for UAVs in a 3D workspace; (ii) accounting for the credit assignment problem via a counterfactual baseline improves planning performance; and (iii) our approach outperforms non-learning-based state-of-the-art approaches in terms of planning performance across varying team sizes and communication constraints. Moreover, we demonstrate our approach applied to a real-world surface temperature monitoring scenario.

A. Experimental Setup

For each experiment, we train for $1e^6$ environment interactions. Then, we execute 50 terrain monitoring missions with changing regions of interest as described in Sec. IV-D. The terrains are of size $50\text{ m} \times 50\text{ m}$ with a map resolution of $r_M = 10\text{ cm}$. We set the planning resolution to $r_P = 5\text{ m}$, bound altitudes between 5 m and 15 m, and use camera field of views of 60° , so that adjacent measurements do not overlap when taken from the lowest altitude. To account for increased sensor noise at higher altitudes, we simulate $p(z_i^t | \mathcal{M}_i^t, \mathbf{p}_i^t)$ to be $\{0.99, 0.735, 0.625\}$ at

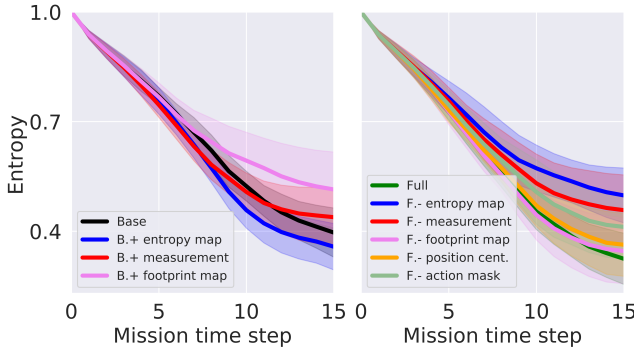


Fig. 5: Feature input ablations. We systematically add (left) and remove (right) components of our approach and show the map entropy reduction over the mission time. The results show that the entropy map input has the largest impact, while the footprint map only contributes to a better performance at the end of a mission. Our full proposed network input (green) performs best.

{5, 10, 15} m altitude. The UAV team consists of 4 agents and the communication radius is limited to 25 m unless reported otherwise. As metrics, we use the entropy of the map state $H(\mathcal{M} | z^{0:t}, \mathbf{p}^{0:t})$ to assess the map uncertainty and the F1-score between the map state $\mathcal{M}$ and the ground truth map to evaluate the correctness of $\mathcal{M}$. We report the mean and standard deviation of these performance metrics in ground truth regions of interest given $B = 15$ measurements.

B. Ablation Study of Network Feature Design

In this section, we provide ablations to show that our proposed network input representation is effective for multi-agent IPP for UAVs in a 3D workspace. Our ablation studies systematically add and remove single input feature maps to assess their effect on the overall planning performance.

Fig. 5 shows the mean map entropy reduction and standard deviation during the mission with varying network input features. Steeper falling curves indicate better planning performance. In Fig. 5-Left, we add single input feature maps to *base* input features (a)-(d), as described in Sec. IV-C, which are necessary to model the IPP problem. In Fig. 5-Right, we remove single input feature maps.

As expected, the proposed full input (green) performs best, i.e. leads to the lowest final map entropy. The entropy map feature (blue) is most beneficial for planning performance, suggesting that it holds most relevant information for learning informative paths. Interestingly, adding the footprint map (pink) slightly harms planning performance, but leads to worse final performance when removed from the input. This indicates that, when combined with the other inputs, the footprint map has a larger contribution at the end of the mission. Intuitively, this is because identifying the currently observed map cells, which is essential for planning the next measurement positions, becomes harder as more measurements are mapped. Adding the current measurement's entropy (red) decreases uncertainty faster during the first ~ 12 mission time steps but harms final performance, presumably since this local feature causes myopic planning bias. As part of the full input with enough global features, however,

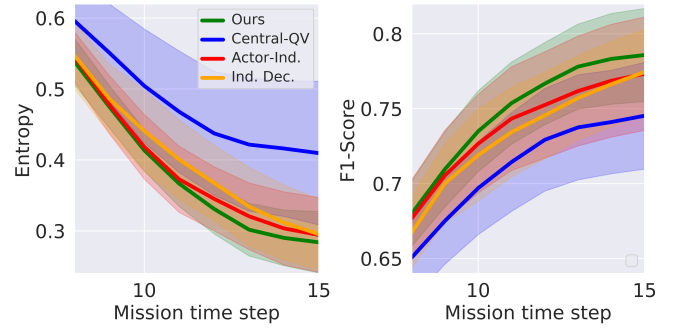


Fig. 6: Credit assignment study. We compare our COMA-approach (green) with variants that do not explicitly tackle the credit assignment problem in terms of map entropy (left) and F1-score (right). Later in mission when the environment is largely explored, our approach enables better cooperation for adaptive mapping.

this feature is significant for planning performance providing additional map information in close proximity. Additionally, not masking invalid actions (light green) and not centring the agents' position map (orange) impairs performance since this removes valuable spatial information for planning. In sum, the results confirm that our network input features ('Full') are most effective for multi-agent IPP for UAVs in a 3D workspace, leading to superior planning performance of our proposed RL-based approach compared to possible variants.

C. Credit Assignment Mechanism Study

The experiments in this section show that accounting for the credit assignment problem results in improved planning performance, which verifies the need for explicit credit assignment mechanisms in cooperative RL-based multi-agent IPP. We perform a systematic study comparing the UAV team's IPP performance with varying advantage functions in Eq. (6) and thus changing policy gradient updates in Eq. (7). To this end, we compare our proposed approach (Sec. IV-B) against three variants of itself, as described in the following. **Central-QV.** Similar to Foerster et al. [20], we verify the effectiveness of the counterfactual baseline in our adaptive IPP scenario by replacing the baseline in Eq. (6) with a state value $V_\pi(s^t)$. This way, we do not estimate an individual agent's contribution but consider the team's joint performance alone. We train two critic networks estimating Q and V , both as described in Sec. IV-D, and change the advantage function to $A_i^t(s^t, \mathbf{u}^t) = Q_\pi(s^t, \mathbf{u}^t) - V_\pi(s^t)$.

Actor-Independent. As described in Sec. IV-B, we utilise a centralised critic exploiting global state information during training. However, to investigate the effect of reasoning about the other team members' actions, we now do not account for them and exclude them from the critic network input. We adapt the advantage function to depend solely on the agent's own action: $A_i^t(s^t, \mathbf{u}_i^t) = Q_\pi(s^t, \mathbf{u}_i^t) - \sum_{u_i'^t \in U} \pi(u_i'^t | \omega_i^t) Q_\pi(s^t, \mathbf{u}_i'^t)$.

Decentralised. In this variant, we remove all global information and consider a purely decentralised critic based on local agent information ω_i^t only. As for the *actor-independent* variant, the critic ignores the other agents' actions using the own agent's state value as a baseline. The

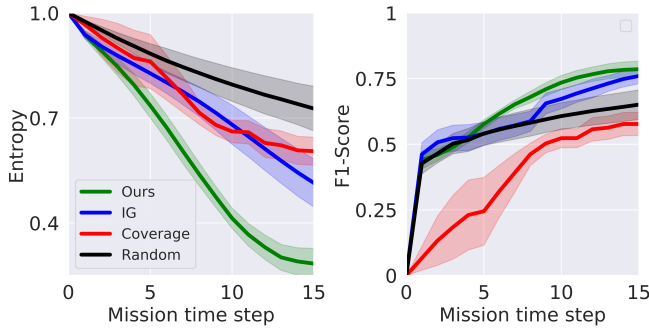


Fig. 7: Comparison of planning methods for 4 agents and a communication range of 25m. By planning paths cooperatively, our approach reduces map uncertainty (left) and improves map accuracy (right) quickest, indicating its superior performance.

advantage function reduces to: $A_i^t(\omega_i^t, u_i^t) = Q_\pi(\omega_i^t, u_i^t) - \sum_{u_i'^t \in U} \pi(u_i'^t | \omega_i^t) Q_\pi(\omega_i^t, u_i'^t)$.

Fig. 6 shows the planning performance using different advantage functions. We focus on performance at later stages of the mission, when most of the environment is explored and cooperation is crucial for adaptive mapping. The counterfactual baseline utilised in our approach (green) performs best, indicating its effectiveness for achieving cooperative behaviour. All variants lack explicit credit assignment mechanisms, which adversely impacts planning performance irrespective of using centralised or decentralised critics. This confirms that our COMA-based algorithm improves planning performance by addressing the credit assignment problem.

D. Comparison against Non-Learning-based Approaches

The following experiments back up our claim that our RL-based approach outperforms state-of-the-art non-learning-based multi-agent IPP approaches. We compare our approach against three methods: (i) an adaptive information gain approach for UAV swarms proposed by Carbone et al. [2], which selects a UAV's action greedily by maximising the estimated map entropy reduction without explicit credit assignment ('IG'); (ii) a non-adaptive coverage lawnmower-like pattern with equidistant (5 m) measurement positions at the best-performing altitude ('Coverage'); (iii) non-adaptive random exploration sampling UAV actions uniformly at random ('Random'). All approaches share the same state space, action space, and action masking strategy.

Fig. 7 shows the evaluation metrics over the mission time for the considered approaches. The information gain-based strategy (blue) and our RL-based approach (green) outperform the non-adaptive methods as they can actively focus on regions of interest. The superior performance of our approach confirms the benefits of addressing the credit assignment problem and the applicability of our learning-based planning to varying terrains.

Next, we show the generalisability of our learning-based approach to different team sizes and communication settings. We deploy our actor trained on the 4-agent setting with communication limited to 25 m in scenarios with 2, 4, and 8 agents, and communication radii of 0 m, 25 m, and unlimited communication without re-training. Table I

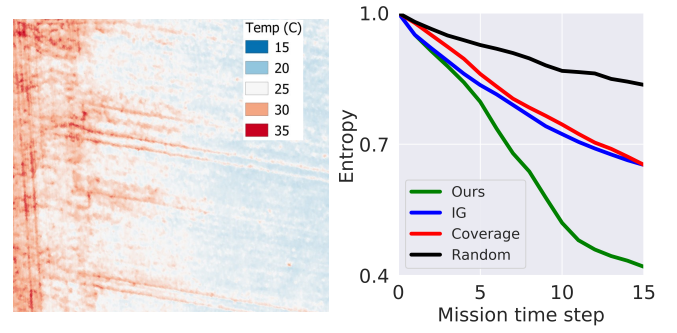


Fig. 8: Real-world evaluation. We deploy our approach on surface temperature data (left) and compare the map entropy reduction of the considered planning methods. Warmer regions are considered as interesting. Although our approach is trained on synthetic data, it outperforms all other methods. Fig. 1 shows the resulting UAV paths planned in this experiment.

shows the planning performance of our approach compared to the information gain-based and the coverage method. As expected, more agents lead to faster mapping, i.e. entropy decrease and F1-score increase. Though changing communication radii lead to small performance drops, our approach consistently outperforms both. This verifies that our RL approach generalises to varying numbers of agents and communication requirements without re-training, showcasing its broad applicability without additional training costs.

E. Temperature Mapping Scenario

We demonstrate the performance of our approach in a surface temperature monitoring scenario using real-world data of a 40 m $\times$ 40 m crop field near Jülich, Germany. The field data was collected with a DJI Matrice 600 UAV carrying a Vue Pro R 640 thermal sensor and is shown in Fig. 8-Left. We discretise the terrain into a 500 $\times$ 500 grid map with resolution $r_M = 8$ cm. The planning grid resolution is $r_P = 4$ m to guarantee the same network input dimensions as used for training. Regions with surface temperature $\geq 25^\circ\text{C}$ are considered as being interesting for adaptive hotspot mapping. Fig. 8-Right shows the map entropy reduction over $B = 15$ measurements for 4 agents and a communication radius of 25 m, using our approach compared to the methods introduced in Sec. V-D. Note that our approach is trained solely in simulation on synthetic data as described in Sec. IV-D. Our approach clearly outperforms all other methods showcasing its applicability for real-world terrain monitoring missions.

VI. CONCLUSIONS AND FUTURE WORK

In this paper, we introduced a novel multi-agent deep RL-based IPP approach for adaptive terrain monitoring using UAV teams. Our method features new network representations and exploits a counterfactual baseline to address the credit assignment problem for cooperative UAV path planning in a 3D workspace. This allows us to successfully outperform state-of-the-art non-learning-based approaches in terms of monitoring efficiency while generalising to different mission settings without re-training. Experiments

TABLE I: Robustness to varying team sizes and communication constraints. We report the mean and standard deviation of entropy and F1-score over 50 trials after 33%, 66%, and 100% of the mission time. Best results in bold.

Setting	Approach	33% Entropy ↓	67% Entropy ↓	100% Entropy ↓	33% F1 ↑	67% F1 ↑	100% F1 ↑
2 agents	Ours	0.8826 ± 0.0267	0.7343 ± 0.0548	0.5802 ± 0.0451	0.3608 ± 0.0661	0.5035 ± 0.0768	0.6268 ± 0.0436
	IG	0.9155 ± 0.0154	0.8449 ± 0.0410	0.7432 ± 0.0801	0.3275 ± 0.0386	0.4353 ± 0.0539	0.5340 ± 0.0857
	Coverage	0.9163 ± 0.0472	0.7845 ± 0.0628	0.6970 ± 0.0332	0.1603 ± 0.0871	0.3687 ± 0.0912	0.4863 ± 0.0418
4 agents	Ours	0.7360 ± 0.0349	0.4137 ± 0.0294	0.2842 ± 0.0408	0.5765 ± 0.0168	0.7350 ± 0.0268	0.7858 ± 0.0308
	IG	0.8302 ± 0.0298	0.6805 ± 0.0554	0.5176 ± 0.0700	0.5396 ± 0.0435	0.6728 ± 0.0375	0.7599 ± 0.0289
	Coverage	0.8615 ± 0.0762	0.6614 ± 0.0317	0.6052 ± 0.0396	0.1603 ± 0.0870	0.3687 ± 0.0913	0.4864 ± 0.0418
8 agents	Ours	0.5621 ± 0.0235	0.2952 ± 0.0403	0.2209 ± 0.0459	0.7115 ± 0.0336	0.7859 ± 0.0322	0.8171 ± 0.0283
	IG	0.6887 ± 0.0234	0.4886 ± 0.0460	0.3077 ± 0.0446	0.6957 ± 0.0260	0.7881 ± 0.0194	0.8576 ± 0.0151
	Coverage	0.8185 ± 0.0832	0.6072 ± 0.0274	0.5149 ± 0.0309	0.2455 ± 0.1281	0.5246 ± 0.0358	0.5800 ± 0.0433
Zero communication	Ours	0.7638 ± 0.0408	0.5445 ± 0.0674	0.3699 ± 0.0486	0.5476 ± 0.0256	0.6846 ± 0.0327	0.7620 ± 0.0282
	IG	0.8284 ± 0.0283	0.6814 ± 0.0529	0.5286 ± 0.0775	0.5396 ± 0.0453	0.6510 ± 0.0442	0.7151 ± 0.0432
Limited communication	Ours	0.7360 ± 0.0349	0.4137 ± 0.0294	0.2842 ± 0.0408	0.5765 ± 0.0168	0.7350 ± 0.0268	0.7858 ± 0.0308
	IG	0.8302 ± 0.0298	0.6805 ± 0.0554	0.5176 ± 0.0700	0.5396 ± 0.0435	0.6728 ± 0.0375	0.7599 ± 0.0289
Full communication	Ours	0.7428 ± 0.0276	0.4623 ± 0.0602	0.3674 ± 0.0843	0.5818 ± 0.0294	0.7077 ± 0.0481	0.7524 ± 0.0465
	IG	0.8294 ± 0.0290	0.6809 ± 0.0501	0.5266 ± 0.0593	0.5389 ± 0.0449	0.6721 ± 0.0432	0.7606 ± 0.0306

using UAV-acquired thermal data validate the real-world applicability of our approach. Future work will investigate heterogeneous robot teams and dynamically growing maps in environments of unknown bounds.

REFERENCES

- [1] D. Albani, T. Manoni, D. Nardi, and V. Trianni, "Dynamic UAV Swarm Deployment for Non-Uniform Coverage," in *Proc. of the Intl. Conf. on Autonomous Agents and Multi-Agent Systems (AAMAS)*, 2018.
- [2] C. Carbone, D. Albani, F. Magistri, D. Ognibene, C. Stachniss, G. Kootstra, D. Nardi, and V. Trianni, "Monitoring and mapping of crop fields with UAV swarms based on information gain," in *Distributed Autonomous Robotic Systems*. Springer, 2022, pp. 306–319.
- [3] W. H. Maes and K. Steppe, "Perspectives for Remote Sensing with Unmanned Aerial Vehicles in Precision Agriculture," *Trends in Plant Science*, vol. 24, pp. 152–164, 2019.
- [4] E. Bondi, D. Dey, A. Kapoor, J. Piavis, S. Shah, F. Fang, B. Dilkina, R. Hannaford, A. Iyer, L. Joppa, and M. Tambe, "AirSim-W: A Simulation Environment for Wildlife Conservation with UAVs," in *Proc. of the ACM SIGCAS Conf. on Computing and Sustainable Societies*, 2018.
- [5] M. A. Goodrich, B. S. Morse, D. Gerhardt, J. L. Cooper, M. Quigley, J. A. Adams, and C. Humphrey, "Supporting wilderness search and rescue using a camera-equipped mini UAV," *Journal of Field Robotics (JFR)*, vol. 25, no. 1-2, pp. 89–110, 2008.
- [6] S. Hayat, E. Yanmaz, T. X. Brown, and C. Bettstetter, "Multi-objective UAV path planning for search and rescue," in *Proc. of the IEEE Intl. Conf. on Robotics & Automation (ICRA)*, 2017.
- [7] Z. Yan, N. Jouandeau, and A. A. Cherif, "A Survey and Analysis of Multi-Robot Coordination," *Intl. Journal of Advanced Robotic Systems*, vol. 10, no. 12, p. 399, 2013.
- [8] R. Bähmann, D. Schindler, M. Kamel, R. Siegwart, and J. Nieto, "A decentralized multi-agent unmanned aerial system to search, pick up, and relocate objects," in *Proc. of the Intl. Symposium on Robotic Research (ISRR)*, 2017.
- [9] G. A. Hollinger, B. Englot, F. S. Hover, U. Mitra, and G. S. Sukhatme, "Active planning for underwater inspection and the benefit of adaptivity," *Intl. Journal of Robotics Research (IJRR)*, vol. 32, pp. 3–18, 2013.
- [10] M. Tzes, N. Bousias, E. Chatzipantazis, and G. J. Pappas, "Graph Neural Networks for Multi-Robot Active Information Acquisition," *arXiv preprint arXiv:2209.12091*, 2022.
- [11] J. Rückin, L. Jin, and M. Popović, "Adaptive Informative Path Planning Using Deep Reinforcement Learning for UAV-based Active Sensing," in *Proc. of the IEEE Intl. Conf. on Robotics & Automation (ICRA)*, 2022.
- [12] A. Viseras and R. Garcia, "DeepIG: Multi-robot information gathering with deep reinforcement learning," *IEEE Robotics and Automation Letters (RA-L)*, vol. 4, no. 3, pp. 3059–3066, 2019.
- [13] H. Bayerlein, M. Theile, M. Caccamo, and D. Gesbert, "Multi-UAV path planning for wireless data harvesting with deep reinforcement learning," *IEEE Open Journal of the Communications Society (OJ-COMS)*, vol. 2, pp. 1171–1187, 2021.
- [14] F. Chen, J. D. Martin, Y. Huang, J. Wang, and B. Englot, "Autonomous Exploration Under Uncertainty via Deep Reinforcement Learning on Graphs," in *Proc. of the IEEE/RSJ Intl. Conf. on Intelligent Robots and Systems (IROS)*, 2020.
- [15] T. Fan, P. Long, W. Liu, and J. Pan, "Distributed multi-robot collision avoidance via deep reinforcement learning for navigation in complex scenarios," *Intl. Journal of Robotics Research (IJRR)*, vol. 39, no. 7, pp. 856–892, 2020.
- [16] S. Y. Luis, M. P. Esteve, D. G. Reina, and S. T. Marín, "Deep Reinforcement Learning applied to multi-agent informative path planning in environmental missions."
- [17] Y. Cao, Y. Wang, A. Vashisth, H. Fan, and G. A. Sartoretti, "CAT-NIPP: Context-Aware Attention-based Network for Informative Path Planning," in *Proc. of the Conf. on Robot Learning (CoRL)*.
- [18] M. Lodel, B. Brito, A. Serra-Gómez, L. Ferranti, R. Babuška, and J. Alonso-Mora, "Where to look next: Learning viewpoint recommendations for informative trajectory planning," in *Proc. of the IEEE Intl. Conf. on Robotics & Automation (ICRA)*. IEEE, 2022.
- [19] A. K. Agogino and K. Tumer, "Unifying temporal and structural credit assignment problems," in *Proc. of the Intl. Conf. on Autonomous Agents and Multi-Agent Systems (AAMAS)*, 2004.
- [20] J. Foerster, G. Farquhar, T. Afouras, N. Nardelli, and S. Whiteson, "Counterfactual multi-agent policy gradients," in *Proc. of the Conf. on Advancements of Artificial Intelligence (AAAI)*, 2018.
- [21] M. Popović, T. Vidal-Calleja, G. Hitz, J. J. Chung, I. Sa, R. Siegwart, and J. Nieto, "An informative path planning framework for UAV-based terrain monitoring," *Autonomous Robots*, vol. 44, pp. 889–911, 2020.
- [22] A. Blanchard and T. Sapsis, "Informative path planning for anomaly detection in environment exploration and monitoring," *Ocean Engineering*, vol. 243, p. 110242, 2022.
- [23] G. Best, O. M. Cliff, T. Patten, R. R. Mettu, and R. Fitch, "DecMCTS: Decentralized planning for multi-robot active perception," *Intl. Journal of Robotics Research (IJRR)*, vol. 38, pp. 316–337, 2019.
- [24] Q. Li, F. Gama, A. Ribeiro, and A. Prorok, "Graph Neural Networks for Decentralized Multi-Robot Path Planning," in *Proc. of the IEEE/RSJ Intl. Conf. on Intelligent Robots and Systems (IROS)*, 2020.
- [25] H. Qie, D. Shi, T. Shen, X. Xu, Y. Li, and L. Wang, "Joint Optimization of Multi-UAV Target Assignment and Path Planning Based on Multi-Agent Reinforcement Learning," *IEEE Access*, vol. 7, pp. 146 264–146 272, 2019.
- [26] A. Puente-Castro, D. Rivero, A. Pazos, and E. Fernandez-Blanco, "UAV swarm path planning with reinforcement learning for field prospecting," *Applied Intelligence*, vol. 52, pp. 14 101–14 118, 2022.
- [27] A. Elfes, "Using occupancy grids for mobile robot perception and navigation," *Computer*, vol. 22, no. 6, pp. 46–57, 1989.
- [28] R. S. Sutton and A. G. Barto, *Reinforcement Learning: An Introduction*. MIT press, 2018.