

Understanding Robot Minds: Leveraging Machine Teaching for Transparent Human-Robot Collaboration Across Diverse Groups

Suresh Kumar Jayaraman
Robotics Institute
Carnegie Mellon University
Pittsburgh, USA

Reid Simmons
Robotics Institute
Carnegie Mellon University
Pittsburgh, USA

Aaron Steinfeld
Robotics Institute
Carnegie Mellon University
Pittsburgh, USA

Henny Admoni
Robotics Institute
Carnegie Mellon University
Pittsburgh, USA

Abstract—In this work, we aim to improve transparency and efficacy in human-robot collaboration by developing machine teaching algorithms suitable for groups with varied learning capabilities. While previous approaches focused on tailored approaches for teaching individuals, our method teaches teams with various compositions of diverse learners using team belief representations, to address personalization challenges within groups. We investigate various group teaching strategies, such as focusing on individual beliefs or the group’s collective beliefs, and assess their impact on learning robot policies for different team compositions. Our findings reveal that team belief strategies yield less variation in learning duration and better accommodate diverse teams compared to individual belief strategies, suggesting their suitability in mixed-proficiency settings with limited resources. Conversely, individual belief strategies provide a more uniform knowledge level, particularly effective for homogeneously inexperienced groups. Our study indicates that the teaching strategy’s efficacy is significantly influenced by team composition and learner proficiency, highlighting the importance of real-time assessment of learner proficiency and adapting teaching approaches based on learner proficiency for optimal teaching outcomes.

Index Terms—explainable decision-making, human-robot teams, group machine teaching, adaptive explainability, team modeling

I. INTRODUCTION

Robots are increasingly becoming an integral part in people’s lives, evolving from human assistants to partners. For safe and effective interaction in human-robot collaboration, it is crucial for humans to understand how robots make decisions. Explainable decision-making aims to clarify the underlying decision-making process of the robot to human collaborators.

Explanations are found to be more useful when they are personalized to the individual [1]–[3]. However, this requires the robots to track the individual’s knowledge or understanding of the system and the surroundings to develop explanations that are better suited for them. In prior work, the human is frequently modeled as an inverse reinforcement learner [4] and learns a robot policy from demonstrations of its behavior

This work was supported by the Office of Naval Research award N00014-181-2503.

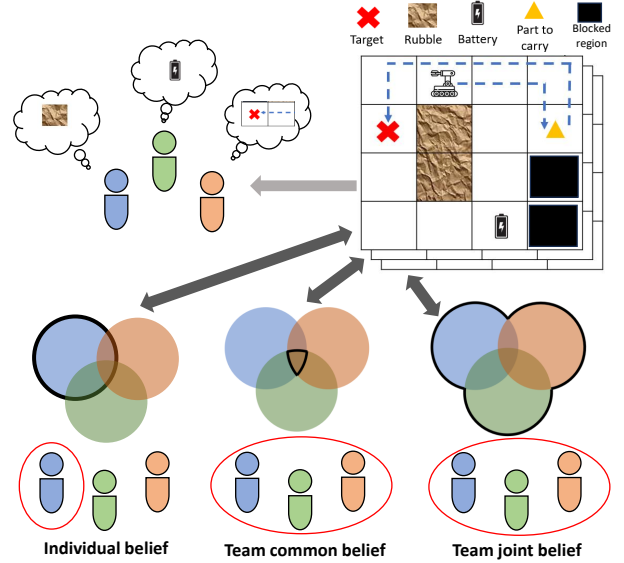


Fig. 1. The figure illustrates the complexity of group machine teaching, highlighting the disparity in understanding from common examples among diverse group members. Personalizing examples to a group is challenging due to varied individual beliefs and learning abilities. Our approach utilizes estimations of individual and collective team beliefs to tailor demonstrations for effective communication of the robot’s policy to the entire group.

[5]. While traditional machine teaching [6] aims to generate informative demonstrations for individual human learners, real-world scenarios often involve the robot working with diverse groups, introducing the need for teaching methodologies that accommodate groups with varying learning abilities, and experiences. This work investigates approaches for robots to effectively communicate their decision-making to groups of human learners through demonstrations.

Teaching a group as a whole instead of individually is preferable, especially with limited time and resources. Consider, for instance, an ad hoc emergency response team tasked with building shelters after an earthquake. They are given a robot that can bring requested items. The robot has limited maneuverability over rubble, a limited range, and may prefer

to recharge when possible. These capabilities and preferences (i.e. its decision-making) must be taught to the team quickly because of the time-sensitive situation. A challenge in group teaching is accommodating individuals with varied learning abilities through a common set of demonstrations. Prior work has shown that it is possible to teach a heterogeneous group of learners using common examples [7], albeit for simple concepts. While groups can also learn from each other through communication and information sharing, we focus only on learning from common examples. Also, group heterogeneity could imply variations in prior knowledge, but we assume similar prior knowledge, focusing solely on differences in the learning ability among the group.

In a classroom, a teacher could personalize the lessons to the class based on various objectives — focus on naive learners who need more support, or on proficient learners, or consider the class as a whole using class average or similar measures — and adapt their teaching accordingly. Yeo et al. [8] explored categorizing learners based on learning rates and provided personalized teaching to each category. Melo and Lopes [9], on the other hand, generated personalized demonstrations for each learner, but at a high teaching cost. An active teacher that personalizes and adapts to the learner can improve learning [10], [11]. But a challenge in groups is identifying to whom the personalization should cater.

Drawing inspiration from pedagogical literature on teaching classrooms [12], our key insight is that *machine teaching can be tailored to a group of learners by considering the group as a whole and generating demonstrations based on the aggregation of the group’s understanding*. In this work, we developed team belief models that facilitate group teaching focusing on the entire team. We utilized a closed-loop teaching framework that effectively incorporates feedback from the robot teacher to aid human learners. We adapted a human belief model to generate simulated human learners of varying learning abilities. We conducted a simulation study to explore how different group teaching strategies affect the group’s learning and how team composition of learners with varying learning abilities — naive and proficient — moderate group learning. Our findings suggests that teaching methods designed for individual beliefs weren’t much affected by how knowledgeable team members were. However, these methods did affect how long team members interacted, depending on team composition. On the other hand, teaching methods that focused on team beliefs helped increase knowledge, especially in groups with more proficient learners.

II. BACKGROUND

Machine teaching for policies: We model the environment as a Markov Decision Process (MDP), given by the tuple $\langle \mathcal{S}, \mathcal{A}, T, R, \gamma, S_I \rangle$, representing the state and action spaces, transition function, reward function, discount factor, and initial state distribution respectively. An optimal trajectory ξ^* is a sequence of (s_i, a, s'_i) tuples obtained by following the robot’s optimal policy π^* . Similar to prior work [13], $R = \mathbf{w}^{*\top} \phi(s, a, s')$ is represented as a weighted linear combination

of reward features. We define a group of MDPs that share $R, \mathcal{A}$, and γ but differ in T_i, S_i , and S_i^0 , as a domain. Sharing the same R ensures that all demonstrations within the domain support inference over a common $\mathbf{w}^*$. We use the MDP formulation to model an *item delivery* task where a robot is tasked with delivering an item in an environment that has rubble, blockages, and a battery recharge station (see Fig.1).

We adapt the machine teaching framework for policies [14] to select a set of demonstrations $\mathcal{D}$ of size n that maximizes the similarity ρ between optimal policy π^* and the policy $\hat{\pi}$ recovered using a computational model $\mathcal{M}$ (e.g., IRL) on $\mathcal{D}$, $\arg \max_{\mathcal{D} \subseteq \Xi} \rho(\hat{\pi}(\mathcal{D}, \mathcal{M}), \pi^*)$ s.t. $|\mathcal{D}| = n$, where Ξ is the set of all demonstrations of π^* in a domain. Once $\mathbf{w}^*$ is approximated through IRL, this approach assumes that the learner can deduce π^* by planning on the underlying MDP. Thus, the objective reduces to selecting demonstrations that are informative in conveying $\mathbf{w}^*$, which can be measured using behavior equivalence classes.

Behavior equivalence class: The *behavioral equivalence class (BEC)* of a policy π is the set of reward functions under which π is optimal. For a reward function that is a weighted linear combination of features, the BEC of a demonstration ξ of π is the intersection of half-spaces [15] formed by the exact IRL equation

$$\text{BEC}(\xi|\pi) := \mathbf{w}^\top (\mu_\pi^{(s,a)} - \mu_\pi^{(s,b)}) \geq 0, \forall (s, a) \in \xi, b \in \mathcal{A}. \quad (1)$$

where $\mu_\pi^{(s,a)} = \mathbb{E} [\sum_{t=0}^{\infty} \gamma^t \phi(s_t) \mid \pi, s_0 = s, a_0 = a]$ is the vector of reward feature counts accrued from taking action a in s , then following π after. Any demonstration can be converted into a set of constraints on $\mathbf{w}$ using (1), with each constraint being a *knowledge component (KC)* [16] that captures a facet of the reward function (e.g., tradeoffs between the underlying reward features). Consider the item delivery domain, which has binary reward features $\phi = [\text{traversed rubble}, \text{battery recharged}, \text{action taken}]$. In practice, we require $\|\mathbf{w}^*\|_2 = 1$ to bypass both the scale invariance of IRL and the degenerate all-zero reward function. If no prior knowledge is assumed, the potential belief space on reward weights would uniformly span the surface of the $n - 1$ sphere (where n is number of domain features) due to the L^2 norm constraint on $\mathbf{w}^*$. We instead assume that learners begin with a prior that action weight is negative (e.g. favoring shortest path, see Fig. 2).

Team modeling: A common way to represent a team characteristic such as team knowledge is by aggregating the knowledge of individuals. Team characteristics are normally represented as average, median, sum, range, minimum, or maximum values of the characteristic of individuals [17]. More recently, team knowledge is represented using a latent *collective intelligence* parameter that is highly correlated with team process and performance [18]. However, operationalizing such a latent parameter is challenging and we thus represent team belief through observable behaviors by aggregating individual beliefs. We focus on two aggregated representations of team belief — common belief and joint belief. We define **common team belief** as the belief that all team members have. It can be

visualized as the intersection of individual beliefs. We define **joint team belief** as the knowledge that at least one individual in the team has, visualized as the union of individual beliefs (see Fig. 3 (b) for visual representations of these).

III. METHODS

In this section, we discuss an approach using particle filters (PF) for modeling individual and team beliefs about the robot’s decision-making, i.e. its reward. We use these different beliefs to select corresponding demonstrations that are shown to the entire team. [11] originally proposed a PF-based approach to model individual human belief that supports iterative Bayesian updates and sampling for generating informative and tailored demonstrations using counterfactual reasoning. We extend this approach to group teaching to model aggregated team beliefs in addition to individual beliefs. We use this model in a closed-loop teaching framework leveraging insights from the education literature and adaptively generating demonstrations based on individual and aggregated team beliefs. In addition, after seeing demonstrations, we provide tests, which collect responses on expected optimal robot trajectory in an unseen environment to evaluate their understanding.

A. Particle filter model of human learner belief

Humans generally perform approximate inference from demonstrations [5] and thus we model the human learner’s belief about the robot’s reward weights using a particle filter, where each particle represents a potential belief about the reward weights [11] and the particle weight represents the belief probability. The particle filter follows a Bayesian update process that uses constraints corresponding to the expected information gain for demonstrations and actual information gain for tests. This formulation enables iterative updates on learner belief from demonstrations and tests.

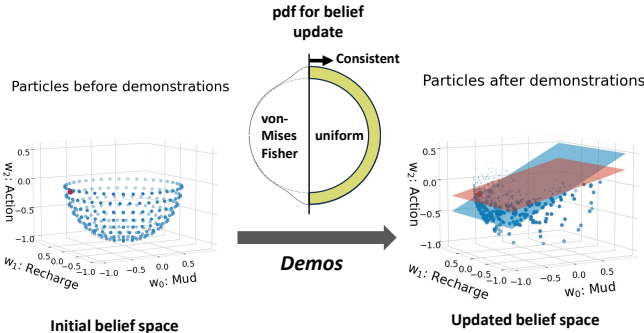


Fig. 2. Update process of a learner’s belief represented by a particle filter. A cross-section of the custom probability density function (pdf) used to update particle weights is shown. Particles consistent with the demonstrated behavior receive higher weights via a uniform distribution (yellow ring), while those on the inconsistent side are weighted less, decreasing exponentially with distance from the constraint, via a von-Mises Fisher distribution. The updated belief distribution is shown on the right.

From demonstrations, constraints on reward weights can be obtained using Eq. 1, by comparing the optimal demonstration with possible counterfactuals. Similarly, the correct test response can be compared with the incorrect learner response to

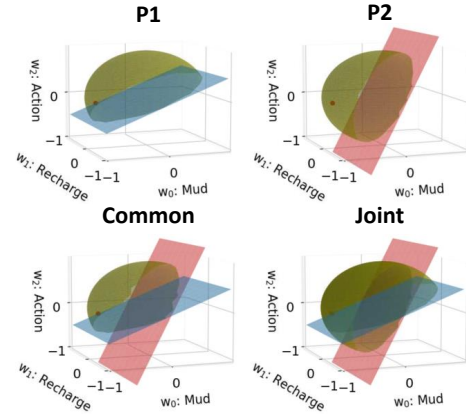


Fig. 3. This figure illustrates an example set of test responses for a team with two individuals, P1 and P2. The test responses are transformed to constraints. The yellow partial spheres show the regions that are consistent with their test response, i.e. agree with the constraint. When their responses are different, the constraints space of their common belief of their tests is their intersection of individual beliefs and that of their joint belief is the union of individual beliefs as depicted. These constraints spaces are used to update the weights of the PF distributions.

get these constraints using Eq. 1. Each constraint c_i generated from the demonstrations and test responses is a half-space constraint, meaning, one side is consistent with the demo or test response and the other is not. Each constraint c_i can be converted to a probability distribution $p(x_i|c_i)$ that is used to update the particle weights w_i of particles x_i .

We use the custom probability distribution (refer Fig. 2) proposed in [11], designed such that any particle on the consistent side of the constraint (yellow region in Fig. 2) is equally valid and could have generated the demonstration (represented by the uniform distribution) and the particles on the inconsistent side are exponentially less likely to have generated the demonstration (represented by the von Mises-Fisher distribution). Fig. 2 shows how the initial learner belief for the delivery robot’s reward gets updated using this custom pdf after seeing the demonstrations. The belief updates after demonstrations and feedback are similar to the *predict* stage and the update based on the learner’s test response is similar to the *correct* stage of Bayesian estimation. The robot teacher uses the learner belief space to sample possible counterfactuals for generating learner-centric demonstrations. To handle potential sample degeneracy in particle filtering, we add a Gaussian noise η when updating the particles [19]. We also add a small Gaussian noise ν while updating the particles after test to account for the teacher’s estimation noise. For more details on the particle filter model refer [11].

B. Team belief modeling

A demonstration’s ability to reveal the underlying reward function via IRL is highly dependent on the counterfactuals considered. Thus the learner’s belief space from which the counterfactuals are sampled critically influences the demonstrations generated. For groups, we model team belief as aggregations of individual member beliefs [17], [20]. The main

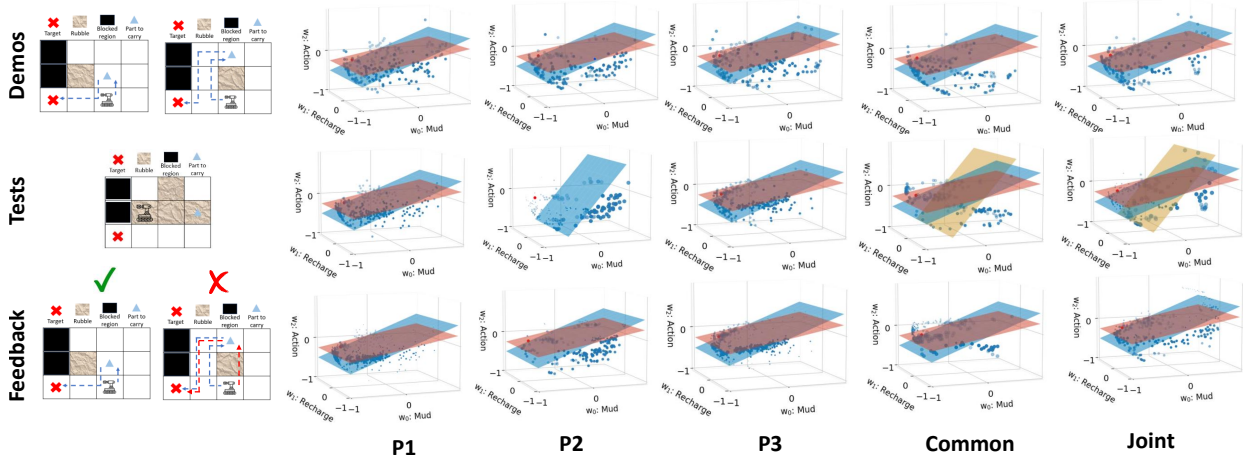


Fig. 4. Interactions and corresponding PF belief updates for a team with three people for the first interaction period. The red particle represents the true reward weight. An interaction period consists of one set of demos related to a KC, followed by a set of tests, and then feedback (corrective or confirmatory). After the demos, all individual and team beliefs are updated based on expected information gain from demos. The distributions are similar since all individuals are expected to learn equally. After the demos, they are provided with test(s) to evaluate their understanding and their responses are used directly for updating individual beliefs and aggregated to update team beliefs. In this case, P1 and P3 got the response correctly and P2 got it incorrect. The difference in the constraint spaces and the updated distributions after tests can be observed for the individual and team beliefs. Confirmatory or corrective feedback is given after the tests and they are expected to learn from either feedback. The distributions are updated to reflect this learning from feedback.

difference in modeling team beliefs is how the particles are updated, specifically how individual constraints are aggregated and used for updating the particle weights.

We envision a group teaching scenario akin to classroom teaching and assume that each team member will have the same kind of interactions, i.e. see the same demonstrations and are provided the same tests. Thus the constraints from the demonstrations will be similar for all individuals. However, let us assume that the team with m members had different responses to a set of n tests, and their constraints are denoted as $C_1 = \{c_1^1, c_1^2, \dots, c_1^n\}$, $C_2 = \{c_2^1, c_2^2, \dots, c_2^n\}$, and $C_m = \{c_m^1, c_m^2, \dots, c_m^n\}$. The update probability for each member is given by, $P_i = \prod_{j=1}^n p(x_i^j | c_i^j)$.

We model common team belief by considering the constraints of all members and representing it as $C_{ck} = \{c_1^1, c_1^2, \dots, c_m^1, c_m^2, \dots, c_m^n, c_1^1, c_1^2, \dots, c_m^n\}$. We assume individuals to be independent. Consequently, the particle filter representing common team belief is updated based on the probability of all aggregated constraints across all tests in the set for all the individuals, given by, $P = \prod_{i=1}^m \prod_{j=1}^n p(x_i^j | c_i^j)$. This aligns with our definition of common belief as the belief that everyone on the team has.

Joint belief, on the other hand, is modeled by considering the set of constraints for all individuals for each test separately and is represented as a set of disjointed subsets, $C_{jk} = \{\{c_1^1, c_1^2, \dots, c_m^1\}, \{c_1^2, c_2^2, \dots, c_m^2\}, \dots, \{c_1^n, c_2^n, \dots, c_m^n\}\}$, where each subset represents the constraints of the individuals. Update probabilities are calculated individually for each team member and the particles are updated based on the maximum probability of any of the individuals, given by, $P = \arg \max_{i \in [1, 2, \dots, m]} \prod_{j=1}^n p(x_i^j | c_i^j)$. This corresponds to our definition of joint belief as the belief that at least one team member has.

Feedback is an effective learning mechanism [21]. Fig. 4 shows the effects of the interactions — demonstrations, tests, and feedback — the team has during one interaction period. An interaction period aims to teach a specific KC, for example, the trade-off between mud cost and step cost for the delivery robot. Each interaction set consists of a set of demonstrations, a set of tests, and a set of feedback (corrective or confirmatory feedback based on whether the response was *incorrect* or *correct*). Confirmatory feedback reinforces the learner’s knowledge while corrective feedback informs them their learning is incorrect and also provides the correct response.

We can see in Fig. 4 that for individuals $P1$ and $P3$ who responded correctly, the belief distribution gets more concentrated within the area of the constraints. This gets further strengthened by the confirmatory feedback they receive. On the other hand, individual $P2$ responded incorrectly indicating they may not have learned the KC. Thus their belief distribution moves away from the constraints region of the KC. However, getting corrective feedback brings the distribution closer to the constraints region. The common team belief moves a little further away because of the incorrect response of $P2$ but is still mostly close to the correct reward and the constraints region.

C. Teaching curriculum development

We employ the methods discussed in [14] to generate demonstrations and tests for teaching the robot’s policy. The approach primarily considers likely learner counterfactuals by estimating the learner’s belief about robot policy (i.e. reward weights) and sampling n possible counterfactuals from this belief space. For every possible robot demonstration in a domain, and for each reward weight, we simulate what the “human” counterfactual to each demonstration would be if

the human had this reward weight in mind and generate the corresponding constraints using Eq. 1 and consolidate all these constraints. We select the demonstration from this set of constraints that maximizes knowledge gain before and after seeing the demonstration. We use the consolidated set of constraints to identify test environments that examine the concept taught in the demos (refer [14] for additional information).

D. Simulated learner model

Learners have different cognitive capabilities and understand at different levels the same information provided. Furthermore, the teaching process is likely to be an adaptive and varied process, catered to the specific learner. Thus our model of the learner should be able to *encompass a wide variety of learner abilities in different teaching contexts*. We again employ a particle filter-based approach to simulate a learner's learning dynamics and belief updates. Similar to the teacher's model, each particle represents a potential belief about the robot's reward function, but they are the learner's self-belief as opposed to the teacher's estimated learner belief. There are two key differences between the teacher and the learner models — first, the learner model updates after seeing the demonstration and feedback only and not after tests since we expect that learners will get information only from the demonstrations and feedback and not from just knowing if they got test responses correct/incorrect [21] and second, the learner model does not have any estimation Gaussian noise (ν) and only the resampling noise (η).

Each individual has a different ability to understand the information conveyed through visual demonstrations. We parametrize this ability as β that can vary for each individual. It is operationalized as the probability mass on the uniform (consistent) side of the custom distribution of the underlying constraint (see Fig. 2). It modifies the information gain in the demonstrations (and feedback) used for particle updates, $IG_{pf} = f(\beta, IG_{demo})$. A higher β implies that learners assign more weights to particles on the consistent side of the constraints, indicating certainty over particles that likely resulted in the demonstrations.

We initialize, $\beta = \beta_0$, as an individual's ability to learn from demonstrations. People get better at learning a concept when they get feedback and are repeatedly exposed to the concept [21]. To incorporate the effects of feedback, we define the feedback dynamics of β as,

$$\beta_t = \beta_{t-1} + \Delta \beta_{t-1}, \text{ with,} \quad (2a)$$

$$\Delta \beta_{t-1} = \begin{cases} \delta \beta_c & \text{if test at } t-1 \text{ is correct} \\ \delta \beta_i & \text{if test at time } t-1 \text{ is incorrect} \end{cases} \quad (2b)$$

where β_c is change in β due to confirmatory feedback, when the learner's test response is correct and β_i is due to corrective feedback, when the learner's test response is incorrect. Corrective and confirmatory feedback have different effects on learning [21]. Fazio et al. [21] found that feedback on incorrect responses led to more learning than feedback on correct responses, particularly in more difficult tasks. Thus we

define $\beta_i > \beta_c$. β_t resets to β_0 , for each new concept as we assume the concepts to be independent. Thus improvement in β due to feedback is contained within the specific concept or KC.

IV. SIMULATION STUDY

We ran a simulation study to evaluate the effects of different teaching strategies on group learning. The strategies differed in the belief space they utilized for sampling possible counterfactuals to generate informative demonstrations. We used $N=8$ counterfactuals in this study. We also examined the effects of the various team compositions of diverse group members on group learning.

A. Metrics

We measure the teaching-learning performance using two measures — (1) the number of interaction periods (N_i) taken to learn the policy, and (2) the average team knowledge level at the end of learning. For each individual, their knowledge level is defined as the proportion of belief space that lies within the BEC region of the robot's policy at the end of the learning session. It is given by $p_{BEC} = \sum_{j=1}^m p_i$, $\forall i \in \epsilon_{BEC}$, where j is the individual, ϵ_{BEC} is the BEC region and i indicates an individual particle and p_i its normalized weight. The team knowledge level, $\overline{p_{BEC}}$ is the average knowledge level of all individuals.

B. Study conditions

The primary condition that we examined was the *teaching strategy* employed to generate the demonstrations and tests. We considered four strategies that samples counterfactuals from four different belief spaces for each interaction period — *individual low*, the belief space is of the individual with the lowest knowledge about the robot's reward, *individual high*, the belief space is of the individual with the highest knowledge about the robot's reward, *common*, the belief space is the common team belief, and *joint*, where the belief space is the joint team belief. The individual, common and joint belief spaces are visualized in Fig. 3. The bigger the belief space, the more diverse the sampled counterfactuals would be. The individuals with lowest or highest knowledge are identified at the end of each interaction period and their corresponding belief spaces are used for the next interaction period corresponding to the strategy employed. We compared these different strategies with a baseline strategy of separately teaching each individual sequentially.

We also examined the influence of *team composition* on the group's learning. Particularly, we considered two categories of learners, *novice* and *proficient*. Novice learners are considered to be beginners and have a low ability to learn from demonstrations, given by a low β_0 . On the other hand, proficient learners have a higher β_0 . For the simulation study, we estimate the distribution of the learner parameters for both the types of learners (see Table I). We considered four team compositions, *all novice*, *majority novice*, *majority proficient*, and *all proficient*.

We expect that group teaching strategies based on group belief, especially the ones based on team belief, target the group as a whole and hence would be able to cater to the entire group resulting in faster group teaching. However, individual teaching strategies, specifically “individual low” is likely to result in higher knowledge as it focuses on the learners who have the least knowledge. So demonstrations catered for them will also improve the knowledge of others. We also expect that teams with more ‘Proficient’ learners will learn quicker than other teams. We expect that the baseline strategy, which teaches each learner separately would take more interactions, would result in better learning because it personally focuses on each learner. Thus we expect, *H1: Group-belief based strategies to have fewer interactions than individual-belief based strategies. H2: Individual low strategy to have the highest knowledge level apart from baseline because of its personalization. H3: Teams with more high learners to have both learn faster and high knowledge levels.*

C. Study setup

We conducted the study for a team with 3 learners. We ran a 5×4 study for the two conditions of *teaching strategy* and *team composition* including the baseline strategy. For each combination, we collected 15 ‘simulated’ teams’ data totaling to 300 teams.

- *Teaching strategy*: baseline, individual low, individual high, common, and joint
- *Team composition*: [N, N, N], [N, N, P], [N, P, P], and [P, P, P], where ‘N’ denotes Novice and ‘P’ denotes Proficient learners.

The learning parameters for each type of learner was estimated (see Table I) from human learner data collected from a user study discussed in [11]. We performed a grid search of the parameters by simulating the teaching interactions from the dataset for all possible grid combinations and choosing the runs that have performance error with $p_{BEC} < \epsilon$

Learner type	$\beta_0 (\mu, \sigma)$	$\delta\beta_c (\sigma)$	$\delta\beta_i (\sigma)$
Novice	0.703 (0.034)	0.033	0.056
Proficient	0.809 (0.025)	0.022	0.052

TABLE I

LEARNER PARAMETERS ESTIMATED FROM HUMAN LEARNER DATA.

The demonstrations are selected based on the KCs and the teaching strategy (which individual/team belief to adapt). The tests are similarly identified to evaluate the constraints conveyed by the demonstrations. For the conditions based on individual beliefs, the individuals with the lowest and highest knowledge are identified at the end of each interaction period, where an interaction period consists of a set of demonstrations, tests, and feedback related to a specific KC. The teaching moves to the next KC only after all individuals in the team learn the current KC. If the team does not learn the KC, the interaction loop continues for the same KC and the demonstrations are generated based on the updated individual or group belief. For aggregated group belief, for example, common belief, it is possible that the responses of teammates are such

that there is no common intersecting constraint region. In such cases, we exclude the conflicting individual(s) constraint spaces and consider the plausible intersecting region formed by most team members.

We developed a closed-loop teaching framework to sequentially generate demonstrations, tests, and feedback to teach and evaluate the team’s understanding of the robot policy. Utilizing scaffolding techniques discussed in [14], we select and sequentially introduce knowledge components (KCs). KCs are broadly defined in the education literature as “a concept, principle, fact, or skill inferred from performance on a set of related tasks” [16]. In our case, KCs represent specific characteristics or distinct constraints of the reward features. For example, the KCs could incrementally teach the bounds on the cost of traveling through rubble given the step cost, followed by bounds on the reward for recharging given the step cost, and then trade-offs between rubble and battery.

After the demonstrations, the group members are given test(s) related to the current KC, to evaluate their understanding. Ideally, human learners will be provided a test environment (see Fig. 4) and asked to map the trajectory they believe the robot will take. For our simulation study, we sample a reward weight based on the individual’s PF distribution and use that as the response to the test environment. The more proportion of particles that lie within the consistent area of the test environment, the more likely the sampled weight vector gets the test response correct. In case they get the response correct, a confirmatory feedback is provided and if they get the response incorrect, a corrective feedback is provided. The β_t of the individual learners are updated according to Eq. 2b.

V. RESULTS AND DISCUSSION

Manipulation check: We wanted to ascertain the distinctiveness of the demonstrations and information provided by the various demonstration strategies. We used the surface area formed by the constraints space (yellow region in Fig. 3) of the demonstrations at each interaction period for comparison, the lower the area, the higher the information conveyed by the demonstrations. Fig. 5 (a) illustrates that the constraints area for the various demonstration strategies are significantly different ($p < 0.01$). Demonstrations based on joint belief result in the most informative demonstrations, whereas those based on the belief of the individual with highest knowledge result in the least informative demonstrations. This is likely because joint belief is a union of individual beliefs, and has a broader distribution resulting in diverse counterfactuals which in turn generate more informative demonstrations.

Experimental conditions results: For the four group teaching strategies, the average number of interactions to learn the reward weights is $N_g = 7.67(1.68)$. Unsurprisingly, the baseline strategy of teaching each learner separately takes more interactions, almost twice as much, $N_b = 17.13(2.32)$. However, contrary to our expectations, the baseline condition had a lower knowledge level, $k_b = 0.59(0.12)$ than the group conditions, $k_g = 0.64(0.21)$, though the difference is not statistically significant.

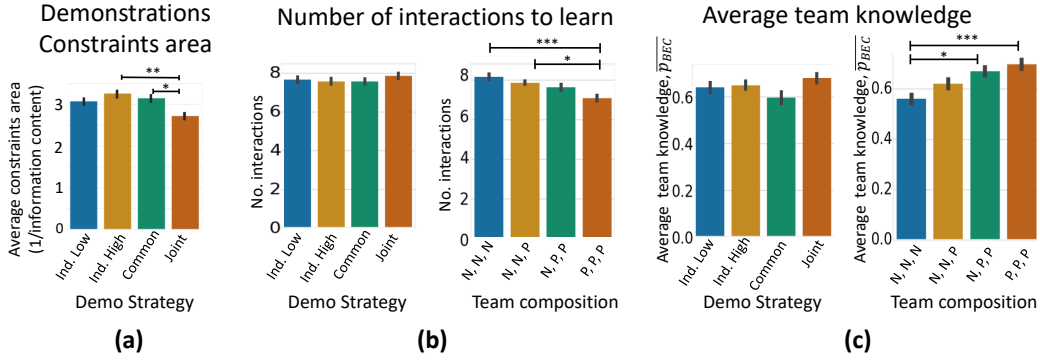


Fig. 5. Experimental results on the effects of demonstration strategy and team composition. (a) All group teaching strategies performed better than the baseline strategy of teaching individuals sequentially in terms of number of interactions. No discernable difference due to strategies for number of interactions. Expected differences in teams, teams with more proficient learners learned quicker, observed. (b) Noticeable differences in average team knowledge observed for strategy but differences are not statistically significant. Also observed that teams with more proficient learners had higher knowledge level.

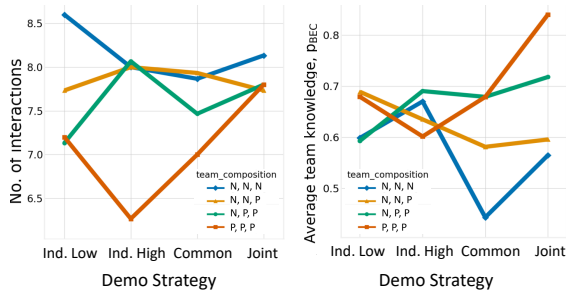


Fig. 6. Interaction effects of demo strategy and team composition on number of interactions and average team knowledge. Differences in variance individual and group belief for both the metrics can be observed.

Since the baseline condition is distinctly different than the other group teaching strategies, we performed two-way ANOVA only on the four group teaching strategies to examine their effects on team learning (see Fig. 5 (b) and (c)). Demonstration strategy did not significantly influence either the number of interactions ($F = 0.42$, $p = 0.74$) or team knowledge level ($F = 1.73$, $p = 0.16$). Team composition, on the other hand, significantly influences both the number of interactions ($F = 4.67$, $p = 0.00$) and team knowledge level ($F = 4.64$, $p = 0.00$), as expected. This is not surprising as with more proficient learners in the team, the team is likely to learn faster and better. We also found a significant interaction effect between demonstration strategy and team composition for team knowledge level ($F = 2.32$, $p = 0.02$).

Fig. 6 shows the interaction trends between demonstration strategy and team composition. While there are no significant interaction effects on the number of interactions, some interesting trends are observed. Individual belief strategies have more variance in learning duration compared to group belief strategies. Group belief strategies, particularly joint belief strategy, is able to accommodate diverse teams and have similar teaching durations. Group belief strategies might therefore be more suitable in situations where the proficiencies of the learners are unclear but still the group has to be taught with limited teaching resources.

On the other hand, we find significant interaction effects for team knowledge level. The interaction trends, are however, interesting as there is less variance in the team's knowledge level for individual belief and more variation for group belief. This could be because of the targeted nature of the individual belief demonstrations, which adaptively samples counterfactuals for upcoming demonstrations from individual belief spaces based on their current knowledge level. This could result in bringing a uniformity of knowledge in the individual belief conditions, particularly for the 'individual low' condition.

To further understand the significant interaction effect we found for team knowledge level, we perform one-way ANOVA for the strategy condition for each team combination and found that demonstration strategy significantly influenced (after applying Bonferroni correction) team knowledge level for homogeneous teams [N,N,N] ($p < 0.05$) and [P,P,P] ($p < 0.05$). Specifically, group belief strategies work well for teams with all proficient learners whereas individual belief strategies work well for teams with all naive learners. For mixed teams, while we do not observe significant differences, we observe similar trends. For teams with more proficient learners perform better with group strategies while teams with more naive learners perform better with individual strategies. The robot can choose the appropriate strategy from the start if proficiency information is available apriori but is more robust if it starts with a strategy and then adaptively change the strategy based on real-time information about the learners proficiency. These nuanced results thus highlight the need to observe and estimate learner proficiency in real-time for more effective teaching.

VI. CONCLUSION

In this study, we aimed to enhance the transparency and efficacy of human-robot collaboration among human groups through explainable robot demonstrations. Our approach involved developing machine teaching algorithms that cater to teams with diverse learning abilities, employing team belief representations aggregated from individual beliefs represented through particle filters. This approach aimed to overcome

the personalization challenges present within heterogeneous groups. Our findings revealed that teaching strategies tailored to group or individual beliefs significantly benefit distinctly different groups characterized by varying levels of learner capabilities. Specifically, we observed that the group belief strategy and joint belief, in particular, was advantageous for groups composed mostly of proficient learners. Individual strategies were better suited for groups with mostly naive learners, though they would take more interactions. We gained deeper insights into the dynamics of group learning, thus laying the foundation for adaptively selecting teaching strategies to facilitate collaborative decision-making in real-time scenarios. However, our study has several limitations. In particular, the simulation did not evaluate the teaching algorithms across multiple domains, nor did it involve actual human learners. Our study had isolated simulated learners and does not consider the nuances of interaction within the group. These limitations highlight the simulation-centric nature of our investigation and suggest the need for empirical validation in real-world settings. Moving forward, we plan to extend this simulation study by conducting a human-subjects user study. This future research will involve actual human learners with diverse learning abilities to assess the efficacy of our teaching strategies. Additionally, we aim to explore the applicability of our methods across various domains and settings to ensure generalizability. By addressing the interaction dynamics within groups, we aspire to refine our teaching algorithms further, ensuring that they are not only effective but also adaptable to the needs of different types of learners and group compositions. By continuing to explore and refine machine teaching approaches, we anticipate contributing to the development of robots that can seamlessly integrate into human teams, enhancing both efficiency and understanding in collaborative tasks.

REFERENCES

- [1] J. Schneider and J. Handali, "Personalized explanation in machine learning: A conceptualization," *arXiv preprint arXiv:1901.00770*, 2019.
- [2] Q. V. Liao, D. Gruen, and S. Miller, "Questioning the ai: informing design practices for explainable ai user experiences," in *Proceedings of the 2020 CHI conference on human factors in computing systems*, 2020, pp. 1–15.
- [3] A. Silva, P. Tambwekar, M. Schrum, and M. Gombolay, "Towards balancing preference and performance through adaptive personalized explainability," in *Proceedings of the 2024 ACM/IEEE International Conference on Human-Robot Interaction*, 2024, pp. 658–668.
- [4] J. Jara-Ettinger, "Theory of mind as inverse reinforcement learning," *Current Opinion in Behavioral Sciences*, vol. 29, pp. 105–110, 2019.
- [5] S. H. Huang, D. Held, P. Abbeel, and A. D. Dragan, "Enabling robots to communicate their objectives," *Autonomous Robots*, 2019.
- [6] X. Zhu, "Machine teaching: An inverse problem to machine learning and an approach toward optimal education," in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 29, no. 1, 2015.
- [7] X. Zhu, J. Liu, and M. Lopes, "No learner left behind: On the complexity of teaching multiple learners simultaneously," in *IJCAI*, 2017, pp. 3588–3594.
- [8] T. Yeo, P. Kamalaruban, A. Singla, A. Merchant, T. Asselborn, L. Faucou, P. Dillenbourg, and V. Cevher, "Iterative classroom teaching," in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 33, no. 01, 2019, pp. 5684–5692.
- [9] F. S. Melo and M. Lopes, "Teaching multiple inverse reinforcement learners," *Frontiers in Artificial Intelligence*, vol. 4, p. 625183, 2021.
- [10] P. Kamalaruban, R. Devidze, V. Cevher, and A. Singla, "Interactive teaching algorithms for inverse reinforcement learning," *arXiv preprint arXiv:1905.11867*, 2019.
- [11] M. S. Lee, H. Admoni, and R. Simmons, "Closed-loop reasoning about counterfactuals to improve policy transparency," in *International Conference on Machine Learning (ICML) Workshop on Counterfactuals in Minds and Machines*, 2023.
- [12] J. M. Carpenter, "Effective teaching methods for large classes," *Journal of Family & Consumer Sciences Education*, vol. 24, no. 2, 2006.
- [13] P. Abbeel and A. Y. Ng, "Apprenticeship learning via inverse reinforcement learning," in *ICML*, 2004.
- [14] M. S. Lee, H. Admoni, and R. Simmons, "Reasoning about counterfactuals to improve human inverse reinforcement learning," in *2022 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*. IEEE, 2022, pp. 9140–9147.
- [15] D. S. Brown and S. Niekum, "Machine teaching for inverse reinforcement learning: Algorithms and applications," in *AAAI*, 2019.
- [16] K. R. Koedinger, A. T. Corbett, and C. Perfetti, "The knowledge-learning-instruction framework: Bridging the science-practice chasm to enhance robust student learning," *Cognitive science*, 2012.
- [17] N. J. Cooke, E. Salas, J. A. Cannon-Bowers, and R. J. Stout, "Measuring team knowledge," *Human factors*, vol. 42, no. 1, pp. 151–173, 2000.
- [18] C. Riedl, Y. J. Kim, P. Gupta, T. W. Malone, and A. W. Woolley, "Quantifying collective intelligence in human groups," *Proceedings of the National Academy of Sciences*, vol. 118, no. 21, p. e2005737118, 2021.
- [19] T. Li, S. Sun, T. P. Sattar, and J. M. Corchado, "Fight sample degeneracy and impoverishment in particle filters: A review of intelligent approaches," *Expert Systems with applications*, 2014.
- [20] S. K. Jayaraman, A. Steinfeld, H. Admoni, and R. Simmons, "Adaptive group machine teaching for human group inverse reinforcement learning," 2023.
- [21] L. K. Fazio, B. J. Huelser, A. Johnson, and E. J. Marsh, "Receiving right/wrong feedback: Consequences for learning," *Memory*, vol. 18, no. 3, pp. 335–350, 2010.