

Scalable Interactive Machine Learning for Future Command and Control

Anna Madison^{*1}, Ellen Novoseller^{*1}, Vinicius G. Goecks^{*1}, Benjamin T. Files¹, Nicholas Waytowich¹, Alfred Yu¹, Vernon J. Lawhern¹, Steven Thurman¹, Christopher Kelshaw², and Kaleb McDowell¹

¹ Humans in Complex Systems, U.S. DEVCOM Army Research Laboratory,
Aberdeen Proving Ground, MD, USA

² U.S. Mission Command Battle Lab, Futures Branch, Ft. Leavenworth, KS, USA

Correspondence: Anna Madison, anna.m.madison2.civ@army.mil

Abstract—Future warfare will require Command and Control (C2) personnel to make decisions at shrinking timescales in complex and potentially ill-defined situations. Given the need for robust decision-making processes and decision-support tools, integration of artificial and human intelligence holds the potential to revolutionize the C2 operations process to ensure adaptability and efficiency in rapidly changing operational environments. We propose to leverage recent promising breakthroughs in interactive machine learning, in which humans can cooperate with machine learning algorithms to guide machine learning algorithm behavior. This paper identifies several gaps in state-of-the-art science and technology that future work should address to extend these approaches to function in complex C2 contexts. In particular, we describe three research focus areas that together, aim to enable scalable interactive machine learning (SIML): 1) developing human-AI interaction algorithms to enable planning in complex, dynamic situations; 2) fostering resilient human-AI teams through optimizing roles, configurations, and trust; and 3) scaling algorithms and human-AI teams for flexibility across a range of potential contexts and situations.

Index Terms—Scalable Interactive Machine Learning, Artificial Intelligence, Human-AI Teaming, Command and Control

I. INTRODUCTION

Future warfare will place unprecedented and dramatically-increased demands on Command and Control (C2)¹ systems.² Maintaining decision advantage over adversaries during multi-domain operations (MDO) will require robust C2 decision-making at shorter timescales in more complex and potentially ill-defined situations. Dispersed, isolated, and mobile command post nodes, forces, and autonomous agents will

need to maintain unity of effort in the face of Denied, Degraded, Intermittent, and Limited (DDIL) communications, while integrating complex data streams into synchronized decision-making capabilities. Yet, current C2 processes are time-consuming and consist of linear sequences of steps, and thus may not adapt well to future battlefield challenges.

To achieve decision dominance on the future battlefield, C2 systems will need to leverage recent breakthroughs in interactive machine learning³ and hybrid human-machine intelligence that indicate the potential for humans and AI to work together to leverage their respective strengths [3]–[5]. There is, however, no one-size-fits-all approach to human-machine integration. Metcalfe, Perelman et al. [6] introduce a *landscape of human-AI partnership*, which posits that the nature of human-AI interaction depends on task complexity, decision timescales, and information certainty. While some tasks may best be solved by AI alone (e.g., humans may not need AI to solve high-certainty, low-complexity, long-timescale tasks), typically, humans and AI agents enjoy complementary, synergistic strengths. For example, AI might efficiently traverse large datasets and explore many possible outcomes, while humans can apply task intuition and common sense and resolve ethical dilemmas, and together, humans and AI may exhibit creativity along different, complementary axes.

AI-enabled decision support systems hold the potential to revolutionize the C2 operations process⁴ to facilitate more robust, efficient, and effective C2 in future operational environments. We propose that research in *Scalable Interactive Machine Learning* (SIML) is critical towards enabling integrated human-AI interaction-based systems that scale to the demands of future C2 (discussed in McDowell et al. [8], ICMCIS 2024 companion paper). First, SIML is critical to streamlining

^{*}These three authors contributed equally to this work.

This research was sponsored by the Army Research Laboratory and accomplished under Cooperative Agreement #W911NF-23-2-0072.

This paper was originally presented at the NATO Science and Technology Organization Symposium (ICMCIS) organized by the Information Systems Technology (IST) Panel, IST-205-RSY – the ICMCIS, held in Koblenz, Germany, 23-24 April 2024.

¹The first “C” in C2, command, is the authoritative act of making decisions and ordering action, with key elements being authority, responsibility, decision-making, and leadership. The second “C” in C2, control, is defined as the act of monitoring and influencing command action through direction, feedback, information, and communications [1].

²C2 systems specify arrangements of people, processes, networks, and command posts [2].

³Methods for humans to cooperate with machine learning algorithms to guide and adapt their behavior toward human-desired intent, e.g., via instructions, demonstrations, or in-the-loop feedback.

⁴The *C2 operations process* is the Army’s framework for organizing and putting C2 into action, and consists of the four major C2 activities performed during operations: planning, preparing, executing, and continuously assessing [7].

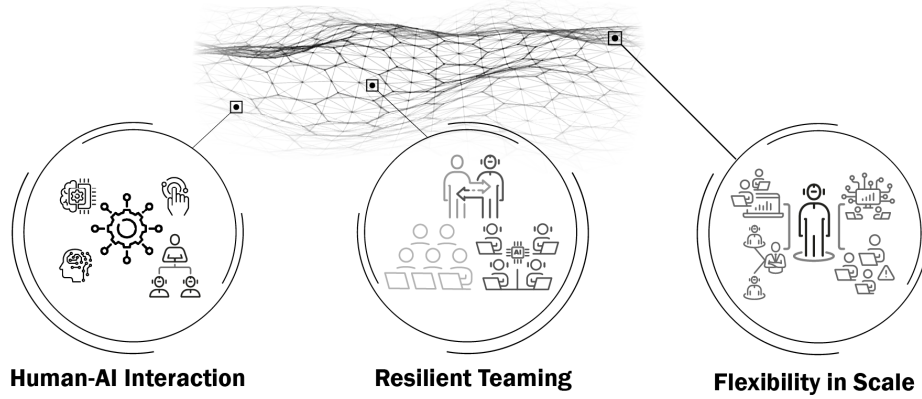


Fig. 1. **Scalable Interactive Machine Learning (SIML) research focus areas.** We propose three research focus areas in SIML to achieve the envisioned C2 developments: 1) developing human-AI interaction algorithms for planning in complex and dynamic environments; 2) fostering resilient human-AI teaming, for instance optimizing team configurations and calibrated trust; and 3) scaling methods developed in the first two research areas to succeed across anticipated future battlefield scenarios.

elements of the C2 operations process, for instance to perform rapid development and analysis of courses of action (COAs), interactively iterating on proposed alternatives based on feedback from C2 personnel; such systems could compress the military decision-making process (MDMP)⁵ or rapidly adjust a COA in real-time based on continuously-assessed battlefield conditions. Second, SIML is critical to maintaining coordination and unity of effort among multiple echelons and across dispersed physical and virtual command post nodes, forces, and AI agents; e.g., SIML models would enhance the ability of dispersed units under DDIL conditions to predict each others' actions in response to unforeseen events. Third, SIML systems could enable AI-based C2 planning to continually adapt based on accumulated data and interactive human feedback, retain knowledge and expertise following personnel rotation, and assist C2 personnel in making optimized decisions based on the most recent and relevant data.

To actualize the envisioned C2 operational developments, this paper proposes three research focus areas in SIML (see Figure 1), covered in the following sections: a) developing interactive machine learning algorithms for robust planning in complex and dynamic environments, b) creating effective and resilient human-AI teams, and c) scaling the approaches in a) and b) to succeed in across a wide range of possible battlefield scenarios. The final discussion also includes practical considerations for operationalizing the envisioned processes.

II. PLANNING IN COMPLEX AND DYNAMIC ENVIRONMENTS WITH INTERACTIVE MACHINE LEARNING

The first research focus area addresses the challenges of jointly leveraging the strengths of both humans and AI to act in complex and dynamic environments. This directly applies to developing well-performing COAs for uncontrolled, uncertain

environments under further challenges such as time pressure, DDIL communications, and unknown opponent intentions.

A. Methods for Humans to Guide Learning and Planning Processes via a Range of Interaction Modalities

Streamlining the MDMP to simultaneously develop and analyze COAs would enable rapid generation and evaluation of many tactical and operational options [9], [10], and would enable adaptation of the MDMP based on commander preferences (note that Section II-B discusses integrating input across multiple people). To develop human-AI interaction methods for such tasks, new research should build on human-guided machine learning techniques that leverage human feedback to train and adapt AI algorithms based on human intentions. Such algorithms leverage humans in a variety of ways, for instance to create knowledge bases (e.g., via instruction manuals or Internet text [11]), to provide instructions (e.g., via language [10] or demonstrations of desired behavior [12], [13]), to give corrections (e.g., allowing users to adjust AI-intended actions [14], [15]), or to provide performance evaluation feedback (e.g., via ratings or relative comparisons identifying user-preferred outcomes [16], [17]). User input could be upfront, as is typical with instructions or demonstrations, or interactively provided during training, as can be more common with corrective or evaluative feedback. Notably, language-based human-AI brainstorming is a promising direction, having been used for text summarization [18]; to train large language models (LLMs) such as ChatGPT, for which human labelers ranked the model's responses from worst to best [11]; and to propose COAs for disaster response scenarios, allowing users to collaborate with an LLM to iterate on plans [15].

Open research problems include how to seamlessly combine different human feedback modalities within a single learning system and how to leverage human feedback to infer human intentions for sophisticated, long-horizon planning tasks such as COA development. These advances would enable C2 person-

⁵In the planning phase of the operations process, the MDMP defines the series of steps used to produce a plan or order [1].

nel to interact with a learning system to adapt proposed COAs and to choose from a variety of interaction types based on situational needs, from demonstrating corrections to catching risky behaviors to communicating shifts in mission objectives to placing constraints that prohibit undesirable actions.

B. Methods for Multiple Humans and Multiple AI Agents to Interact in Complex and Dynamic Environments

In the machine learning research community, several methods have been proposed for decentralized planning with multiple AI agents [19]–[22]; yet, there has been relatively little research on methods for interactive human-guided learning that extend to multi-agent systems or that integrate feedback from many dispersed humans. Relevant to C2, interactive AI methods that plan over multi-agent systems would enable a COA development and analysis system to plan over many distributed forces and assets, which could be either humans or AI agents, while receiving feedback from C2 personnel to improve the proposed operational and tactical recommendations. Meanwhile, interactive AI methods that receive feedback from groups of humans would enable the system to intelligently integrate input from a range of C2 personnel, potentially with different functions, areas of expertise, biases, and mindsets. Such systems could weight feedback from various humans, e.g., depending on their expertise or input quality, and would assist with smooth transitions following personnel rotation. These methods for multi-human multi-AI interaction will lay the foundation for resiliency to DDIL conditions and to other interruptions in human availability, discussed further below.

Existing learning methods for multi-agent systems typically either optimize behavior relative to a pre-specified numerical objective function [19], [23]–[25] or leverage upfront human demonstrations [25]–[28], rather than enabling interactive human feedback throughout learning. Further research is needed to develop methods for humans to intuitively guide large groups of interacting agents, e.g., via paradigms such as multi-agent reinforcement learning (RL).

Most existing human-guided learning methods either learn from a single human or treat feedback from multiple humans as if it came from a single person [29], [30]. Several works introduce methods for crowdsourcing human feedback, that is, obtaining human feedback from an ensemble of humans with different biases and reliability levels. These approaches query humans for discrete classification labels [31]–[34] or pairwise comparisons [35]–[38] or obtain task demonstrations from multiple people [39]. Yet, these approaches often assume that all human data is provided upfront, rather than given interactively [31]–[35], [37]–[39], do not consider sequential decision-making tasks [31]–[34], [37], [38], and/or require that every human label every query [31], [36]. Important open problems in learning from multi-user feedback include, but are not limited to, enabling more and differing interaction modalities across the crowd, integrating real-time human feedback (rather than assuming all human data is available upfront), handling limited or intermittent availability of various humans, determining when (and how) to learn from potentially conflict-

ing feedback across the crowd, effectively propagating critical information from minority crowd members, and accounting for the humans' various strengths, weaknesses, and biases.

C. Human-AI Interaction Methods in which Humans Interact at Multiple Levels in a Hierarchy

AI systems for C2 processes will need to integrate and synchronize multiple Army echelons' actions and effects. We propose to leverage hierarchical AI methods for learning and planning, which not only have the potential to learn faster than flat architectures [40]–[43], but also are built to decompose tasks for multiple levels of planning and execution [42], [44]–[53]. Algorithms for human-AI interaction can especially benefit from a hierarchical structure, since this aligns with how humans naturally process complex tasks, breaking them down into simpler, manageable subtasks [54]. We expect AI systems that can efficiently decompose tasks into a hierarchical structure to more effectively learn from humans and to decompose high-level goals into low-level actions. Such methods hold the potential to enable future AI systems for cross-echelon C2 that function across the C2 operations process.

Further research is needed to develop hierarchical learning and planning methods applicable to complex and dynamic environments, in which humans interact at multiple levels in a hierarchy. While existing hierarchical learning methods have shown promise, these methods typically only study two-level hierarchies, which are much simpler than hierarchies in real-world applications. In addition, while some works study human feedback in hierarchical learning settings, e.g., via language and preferences [51]–[53], [55]–[58], the human feedback is only given at a single level in the hierarchy. Further work should also develop methods to facilitate proper synergy and communication between networks of humans and AI agents organized in multiple hierarchical levels [59].

D. Human-AI Interaction Techniques Robust to Limited Communication and Imperfect Information Scenarios

Previous research has proposed methods for multi-agent planning with limited inter-agent communication [60]–[63]; however, these methods consider information sharing in constrained forms, e.g., agents can only share specific predictions [61], [62], [64], the agents' ability to communicate is determined solely based on geographic proximity [60], [64], or information transmission is subject to time delays [63]. Further research is needed to develop algorithms that are robust to more sophisticated constraints on communication, analogous to DDIL conditions encountered in warfare. Another open problem involves modeling the intentions and future actions of other agents in complex and dynamic situations. In contrast, existing works [24], [25], [62], [65] do not robustly operate under partial observability or reliably identify decision-making factors or model-changing behaviors of other agents [66].

The proposed research will enable AI systems for C2 that help integrate and synchronize dispersed forces and MDO effects across dispersed command post nodes to operate under DDIL conditions, as well as with the incomplete knowledge of

other entities' plans that might arise under DDIL conditions. The problem of maximizing synergy among cooperating AI agents [59] (see Section II-B) becomes even more critical under DDIL conditions, as agents observe differing information and experience hindered communication capacity. Estimating teammates' actions under DDIL is challenging with or without AI; however, we contend that AI has the potential to improve plan intricacy and to determine plan feasibility. Further research should develop methods that model the intentions, beliefs, capabilities, and future actions of other agents involved in a complex and dynamic scenario. Techniques such as active sensing could identify key pieces of information to send or observe to reduce uncertainty in the future actions distribution.

E. Human-AI Interaction Techniques for Leveraging Adapting Databases of Human-Generated Data

Recent advances in natural language processing [67]–[70] have demonstrated how machine learning algorithms can learn from widely available human-generated text corpora. People can further specialize these models to solve desired language-based tasks via *in-context learning* [71]–[73], in which users simply adapt models via short snippets of text or information (*prompts*), as opposed to training them from scratch with task-dependent text corpora [74]. These models can be leveraged to generate plans [15], to reason about how to complete an objective in terms of subtasks [51], [52], or to propose new skills or behaviors to maximize a language-defined objective [53], [75]. However, there remain a number of open research questions, including how to provide tailored information to humans based on their roles, their functions, or the context in which they are performing. Another open question involves how to design AI agents that quickly adapt to new information when learning from large corpora of human-generated data, particularly when these databases expand over time to include new human data with multiple potentially-shifting human mindsets and biases. Such methods will allow C2 planning algorithms to leverage large databases of multi-modal data—from annotated maps to recordings of COA development sessions to written after action reviews—and to adapt quickly as these databases grow, e.g., maintaining robustness following personnel rotation.

III. EFFECTIVE AND RESILIENT HUMAN-AI TEAMS FOR COMPLEX, DYNAMIC TASKS

The seamless interaction of human-AI partnerships in any complex, dynamic environment requires understanding how to create and optimize resilient teams. The second research area focuses on the gaps related to understanding how to define the roles of humans and AI and the configurations of human-machine teams, selecting and training optimal human partners, and facilitating trust and communication across human-AI teams. By solving open research problems related to this focus area, SIML-enabled systems can be better designed for and integrated with C2 personnel performing future C2 processes.

A. Defining Resilient Human-AI Roles and Configurations

The emergence of SIML tools will shift roles of C2 personnel, and novel team configurations must be determined to com-

bine human and AI roles together to efficiently and effectively accomplish C2 tasks. Roles and configurations within human teams are typically allocated based on the knowledge, skills, and abilities of each team member, but SIML tools enable non-traditional, flexible team roles and configurations [76]. Machine intelligence will be better suited for some roles and tasks that are, for example, data processing heavy or repetitive, freeing up humans for new roles, for instance, ensuring the reliability of data and recommendations from AI systems. Within the human-AI team, each partner should be able to detect limitations in the system and influence tasking to overcome shortcomings in the system by changing the configuration or roles, e.g., via task reallocation, which can be accomplished through creating team design patterns (TDPs) [77]. TDPs are a promising approach to create efficient human-machine teaming workflows, and could be used to modernize processes and procedures of functional and integrating cell members within a headquarters staff. TDPs provide generic reusable behaviors across team members for supporting effective and resilient teamwork that allow for a range of combinations of human-machine partnerships based on task and type of work (i.e., cognitive vs. physical) [78], [79]. Since a TDP breaks down a task into sub-processes based on capabilities of human or AI systems [80], this approach provides a design solution to address meaningful human control or ensure humans are able to make informed, timely decisions over AI systems to achieve (or prevent) a given outcome [81].

The structure of a human-AI team should be dictated by the types of human-AI interactions leveraged, such as mediating AI decision recommendations or providing feedback on RL policies [76]. Emerging research focused on understanding consistent roles within human-AI teams suggests supervising, collaborating, and overseeing as potential human roles [82], [83]. An SIML system, for example, focused on creating an accurate Common Operating Picture (COP) may have data wranglers and human deception checkers to validate incoming data from sensors and sources. Building on this, information from SIML systems should convey tailored information depending on human roles, function, and context. C2 personnel focused on targeting, for example, could distill or represent information within the AI system that allows for optimizing task roles, skill knowledge, and context. Role testing and evaluation will be critical steps in defining human and AI roles, for instance facilitating improvements in aspects of human workload or exploitation of the AI. For example, Ramchurn et al. [84] found that a human-AI partnership performed best on a team scheduling task when a human role required mediating between human and AI partners, allowing the AI to overcome prediction constraints and unexpected outcomes.

Open research problems in this area include understanding how human roles in a multi-human and multi-AI team change as the AI systems become more capable, how roles and compositions need to change as the complexity, time-sensitivity, and uncertainty of the problem space increase, and how to create TDPs efficiently for C2 processes. Operational requirements will drive development of SIML tools and interactions, but

addressing these problems will allow for a range of potential available resources in human and AI partners. Optimization of human and AI roles along with their configuration within the team must stem from systematic testing and evaluation based on comparing overall performance outcomes to determine how to create teams that scale and apply across a multitude of operational settings. Development of C2 personnel and staff TDP libraries, as part of Army knowledge management, could serve as doctrinal templates that help form foundational planning standard procedures and guide collective C2 work [78].

B. Effective Human-AI Partnerships through Selecting and Training Optimal Human Partners

A budding area of research seeks to understand characteristics of human partners that result in better human-AI integration. Specifically, how do cognitive abilities, personality traits, and attitudes influence interactions with SIML systems, and what makes some people optimal human partners for a given AI system? A theoretical framework proposed by Hoff et al. [85] suggests that dispositional characteristics, such as personality traits, gender, and age, along with personal experiences and context of use, can greatly influence dynamics of calibrated trust and subsequent interactions with AI or autonomous systems. Achieving calibrated trust results in minimizing distrust that can lead to errors, disuse, and over-reliance, while reducing over-trust that can lead to miscalculations in risk [86], [87]. Training programs aimed at calibrating trust in human-machine teams have proven effective, since trust can evolve over time [88].

A key aspect of human-AI interactions is the formation of a shared mental model, i.e., a shared understanding of the team's task, each team member's role in it, and their respective capabilities, which allows for the predictions of needs and behaviors of other team members [89]. This suggests that traits such as empathy, theory of mind, or emotional intelligence might allow a human to be flexible in their mental model of the AI agent. Selection and understanding of these traits in human partners could then advantageously remedy inherent AI bias, provide over-correction in training data, or result in more flexible mental models that are more resilient to changes in AI systems. Emerging work with human-AI teams in code translation tasks suggests human teammates can overcome limitations in AI capabilities by successfully performing tasks with imperfect AI [90]. Similarly, reasoning over quantified uncertainty is a human skill that can be improved, in some contexts, with deliberate practice [91], [92]. SIML systems will differ in their sophistication and ability to be used effectively amongst the general population. Research is needed to identify characteristics of optimal human partners and human-AI compatibility; such work will aid in alignment of roles and configurations to result in the best performance outcomes and robust, ethical human-AI partnerships [93]. While work is ongoing to define and model AI technological fluency [94], additional predictors of successful human-AI partnering might go beyond more general use of AI. Compositionality is a concept

related to fluid intelligence that explains cognitive flexibility as the ability to tackle complex tasks by combining simpler cognitive operations. Seeking this ability in both human and AI teammates could serve as a foundation for successful human-AI partnerships [95]. Given that AI capabilities and vulnerabilities change rapidly, and these changes are likely to accelerate, human ability to promptly identify and adapt to changing operational characteristics may be paramount [96].

Open questions in this area include better understanding what human characteristics are relevant to successful interaction with SIML systems; how best to develop a workforce with those characteristics, be it through selection, education, training, or other approaches; and how to make sure humans are able to keep up with the rapidly changing AI landscape. Addressing these questions would potentially enable a greater population to effectively modify and guide AI systems, and ensure humans thrive as members of human-AI teams.

C. Communication through Explainable AI and Shared Mental Models

Within teaming situations, communication [93] and development of a shared mental model are important for humans to accurately predict behaviors of both AI [89] and other humans. Just as a human teammate might be required to provide supporting reasoning and evidence for a recommendation, AI will need to provide explanations to ensure a shared understanding of the costs and benefits of a particular recommendation. While AI decision-making methods such as RL are capable of rapidly exploring numerous decision sequences through policies that map observations to actions, these policies often result in black-box functions that are challenging for humans to explain or justify due to their unintuitive nature [97].

Explainable systems [98] may help humans to build calibrated trust [99], [100] in the AI system, but explanations that are not well suited to the task or context might lead to miscalibrated trust [101], [102]. Miller [103] argues that AI must leverage humans' mental models of the world to generate effective explanations. Consistent with this idea, Neerincx and colleagues [104] propose an explainable AI approach that is transferable across domains and focuses on explaining an agent's behavior in terms of its data-driven (perceptual; e.g., contrastive explanations and confidence metrics) and model-based (cognitive, e.g., beliefs and goals) processes while considering the use context and human user capabilities in explanations. This is consistent with another approach to increasing trust between humans and AI through communication by accounting for uncertainty of AI system recommendations and solutions via AI-provided confidence ratings [100]. Approaches for explainability can be integrated and executed through use of TDPs [81], [103], which may shape the human partner's mental model of the AI [77], [103]. But, formation and shaping of the AI mental model is not so easy. One approach is to create an AI compatible representation of human behavior or knowledge and for the AI to validate its understanding back to the human user. Recently, this approach was successfully applied to mission planning,

specifically understanding commander’s intent, to accelerate planning and decision making processes [10].

Open problems in this area include creating effective methods for increasing communication and shaping shared mental models across human-AI partnerships and that apply to complex, high-dimensional, and dynamic planning tasks that will be found in future C2 operational environments. Research is needed to extend policy-based human-AI interaction approaches to predict how actions will unfold [105], [106] while accounting for uncertainty. Recent progress has been made in conveying uncertainty in relatively simple contexts [107], [108], but these methods become computationally intractable under high complexity [109] or account only for model uncertainty [110], and thus modeling uncertainty remains an open problem. Further research should also develop methods that convert policies to explainable plans, e.g., via language, annotated maps, or trees of potential outcomes. Such approaches could leverage generative AI tools such as LLMs and image generation models [11], [68], [111], [112] and combine machine learning methods for sequential planning (e.g., RL) with generative AI [25] to convey information to human users that improves a shared mental model. Related, a remaining challenge is measuring and eliciting shared mental models in experimentation to optimize human-AI partnerships [89].

IV. FLEXIBLE AND SCALABLE DECISION-MAKING WITH INTERACTIVE MACHINE LEARNING

To deploy human-AI interaction systems not only in simulations and simplified exercises, but also in real-world situations, AI algorithms will need to scale flexibly and robustly in response to a variety of dynamic factors. Our third research focus area considers flexibility across scales, particularly in regard to decision-making timescales, the numbers of interacting human and AI agents, the type of decision-making hierarchy, and the problem sphere or scope. Critically, AI systems must not only scale, e.g., as the number of humans increases, but must also flexibly adapt as scales vary bidirectionally in real-time, e.g., if suddenly only two humans are available to interact with the AI system, and as a result, certain expertise is missing. We discuss open research opportunities in each of these areas. Addressing these gaps will be key to enabling scalable human-AI partnerships that can learn and plan in complex and dynamic environments, as are prevalent in C2.

A. Temporal Scalability

Algorithms for human-AI interaction must function smoothly across both longer and shorter decision-making timescales. This is important for C2 applications, since depending on circumstances, the available decision-making time could vary from minutes or hours to days, while possibly respecting further constraints such as the 1/3-2/3 rule.⁶ Algorithms should adapt their behavior given the available time,

⁶From U.S. Army Doctrine Publication 5.0, “The Operations Process”: “Commanders follow the ‘one-third, two-thirds rule’ as a guide to allocate time available. They use one-third of the time available before execution for their own planning and allocate the remaining two-thirds of the time available before execution to their subordinates for planning and preparation” [7].

e.g., by intelligently trading off between factors such as solution optimality, the number of explored decision alternatives, and a temporal budget. Furthermore, temporal budgets could in turn constrain resources such as computational time and human time spent interacting with the AI. Such trade-offs could affect AI behavior in many ways, for instance by informing the number of AI-generated alternatives (e.g., COAs) shown to C2 personnel, the number of simulations or wargames in which these alternatives are evaluated (e.g., COA analysis), the choice of *which* alternatives to present to personnel (e.g., COAs with more predictable outcomes vs. high-risk, high-reward alternatives), and the chosen modalities of human-AI interaction (e.g., querying personnel for detailed feedback vs. a simple COA selection among presented alternatives).

While some existing AI decision-making methods consider a query budget that limits algorithm exploration [113]–[115], consider the time required to process input features [116], trade off between speed and accuracy [117], or use a temporal budget to inform the fidelity at which to run simulations (i.e., balancing between simulation speed and accuracy) [118], these methods apply to limited settings that are significantly simpler than C2 scenarios. Furthermore, they often utilize no human interaction [113], [115], [116], [118], or at most a single human interaction modality [114], [117], and often only employ the budget constraint to determine when to halt algorithm execution [113]–[115]. Further research should develop human-AI interaction methods that can flexibly scale to a range of temporal constraints, for instance determining the best allocation of human and computational resources and the optimal human interaction approach for the available time. Future research should also consider how multiple decision-making AI agents operating at different timescales can best synergize with one another, as well as with humans.

B. Human-AI Interaction Scalability

Human-AI integration techniques must also scale across changing numbers of human and AI agents in the system, including as availability of human or AI agent resources increases or decreases in real-time. Algorithms should be capable of integrating feedback from many people, while retaining robustness when only a few people—or even just a single person—are available to interact or provide feedback. Future work should develop decision-making methods that not only can accept feedback from multiple humans (Section II-B) and are cognizant of these humans’ diverse roles and expertise (Section III-A), but also that are flexible given shifts in the numbers and expertise of accessible humans, including in dynamic scenarios, to maximize human resources while filling in when people become unavailable. In C2 scenarios, such approaches would integrate feedback from many battlefield assets and C2 personnel, while remaining robust when only few C2 personnel are present, seamlessly adapting based on the specific roles and functions of available staff.

Decision-making algorithms must also scale to robustly direct either many or few AI agents. Existing algorithms for multi-agent learning and planning often become computation-

ally intractable as the number of AI agents increases or cannot adapt if AI agents are dynamically added to or removed from the system [119]–[122]. We envision improved algorithms for human-AI integration that scale to plan over many (hundreds or more) forces and assets on the battlefield, while also reliably generating plans when comparatively few assets are present, or when asset availability shifts (e.g., decreases) in real-time. Ensuring synergy between many AI algorithms that receive input from many humans is an open research area; for instance, Friston et al. [59] suggest research avenues for coordinating among a network of many interacting AI systems to achieve collective intelligence.

C. Hierarchical Scalability

To tackle sophisticated problem scenarios such as arise in C2, human-AI partnerships will need to be organized hierarchically, with structured human-AI interactions occurring at multiple hierarchical levels (see Section II-C). Furthermore, however, these algorithms will need to flexibly operate across a range of potential hierarchical structures, switching between them as needed, e.g., to rapidly adapt to changing battlefield conditions. Further research should develop hierarchical algorithms for human-AI interaction that can scale across many hierarchical organizations, e.g., with varying numbers of hierarchical levels and extending to tree-like hierarchical structures. These algorithms should also flexibly adapt in real-time if hierarchies change based on situational needs. In the C2 space, such algorithms will promote robust human-AI interaction across various multi-echelon organizational structures, assisting each echelon and unit individually while maintaining unity of effort. They could also help to coordinate among hierarchies within different countries with different equipment, data, etc. in a single joint mission command situation, e.g., directly facilitating NATO interoperability [123].

D. Problem Sphere Scalability

In complex domains such as C2, learning and planning algorithms should identify the sphere (or scope) of the problem that is most crucial for effective decision-making, and then autonomously adapt decision-making to the key factors within this sphere. For instance, such a capability could help to focus algorithmic planning for C2 on the most relevant geographic regions; perhaps, for the current operation, activity in the east must be estimated more accurately than activity in the west, while for a different operation, activity must be precisely estimated over a larger region of the battlefield. Meanwhile, AI planning should also scale its focus to concentrate on the most relevant warfighting functions; for example, perhaps for the current mission, fires must be estimated more accurately than sustainment. AI algorithms may need to adapt their focus in real-time to changing regions of interest (whether geographical, operational, etc.), while different echelons or dispersed forces may have differing modeling priorities or require differing levels of abstraction. Algorithms that leverage previous experience and human interaction to scale their computational focus to a problem's most critical aspects

would facilitate improved utilization of constrained temporal and computational resources. While multi-resolution modeling can enable AI to learn from data at different abstraction or granularity levels [124]–[126], future research in human-AI integration should expand on these techniques to promote both scalability and adaptability across a broad range of problem spheres, as may arise in complex military situations.

V. DISCUSSION

An SIML-enabled C2 system has the potential to revolutionize C2 personnel effectiveness and efficiency to achieve decision advantage over peer and near-peer adversaries. Yet, additional research is needed to develop the SIML science and technology capabilities necessary for approaches to be practical and robust in the C2 domain. Here, we proposed and described three critical research focus areas for enabling SIML systems for C2. First, successful SIML relies on leveraging human-AI interactions for complex planning, e.g., for multiple humans and multiple agents to collaborate together via a variety of interaction modalities and across hierarchical levels. We expect research in this area to develop over the next decade. Second, while the current MDMP is a fundamentally human endeavor, an SIML-enabled MDMP must reconsider who and what AI systems comprise a given team, along with optimal team roles and configurations given available C2 personnel. Future research should aim to better understand how to optimize team performance through selection, training, and role assignment. We expect this newly emerging field of human-AI teaming performance and optimization to develop over the next decade, whereas the application of the science may begin to emerge as early as 2040. Third, advances in human-AI integration and resilient teaming will only be meaningful for C2 situations if they are flexible and scalable to a broad range of potential operational contexts. While perhaps the most difficult to execute, this research focus area is critical to the transition of such SIML-enabled systems to real-world C2 contexts. Given the validation and testing requirements of military technology along with additional security standards and protocols, we expect to see smaller applications first appear more widely within the next 5 years, followed by more advanced SIML-enabled systems by 2040.

Finally, practical considerations must be taken into account when deploying SIML science and technology capabilities at the tactical edge. In line with industry trends of incorporating AI technologies into smartphones and personal computers, though in a miniaturized version, we foresee that the computing and power requirements for future AI technologies will be met for warfighters at the tactical edge. This vision aligns with ongoing advancements in AI, suggesting that these critical capabilities will be accessible to future C2 personnel potentially as early as 2040. In addition, there are many open considerations in the implementation of SIML-systems for C2 operational environments, such as the choice between pre-learned policies versus online learning or the type of architecture for a given system. Perhaps most important are practical considerations that impact aspects of interoperability between

NATO and Ally forces. Interoperability allows different Allies to operate together in a coherent and efficient manner to achieve tactical, operational, and strategic objectives through common communication, standards, procedures, equipment, and doctrine. Related decisions regarding data and model management are critical, encompassing how data is stored, accessed, and utilized for learning, as well as how models are maintained, updated, and managed over time. As time evolves, NATO and other multi-national agreements, such as STANAG, a NATO Standardization agreement, will need to adapt to emerging technology, such as SIML-enabled systems, to modernize interoperability objectives [123].

ACKNOWLEDGEMENTS

The authors thank Mr. Lee Vinson from the Apex Analytics Group and the Mission Command Battle Lab at Fort Leavenworth, Kansas, for their insightful feedback and discussion when refining the concepts proposed in this paper.

REFERENCES

- [1] Headquarters, Department of the Army, *Army Field Manual 6-0: Commander and Staff Organization and Operations*, 2022.
- [2] Headquarters, Department of the Army, *Army Doctrine Publication 6-0, Mission Command: Command and Control of Army Forces*, 2019.
- [3] N. Case, "How to become a centaur," *Journal of Design and Science*, 2018.
- [4] P. Scharre, "Centaur warfighting: The false choice of humans vs. automation," *Temp. Int'l & Comp. LJ*, vol. 30, p. 151, 2016.
- [5] R. Zhang, F. Torabi, L. Guan, D. H. Ballard, and P. Stone, "Leveraging human guidance for deep reinforcement learning tasks," in *International Joint Conference on Artificial Intelligence*, 2019.
- [6] J. S. Metcalfe, B. S. Perelman, D. L. Boothe, and K. McDowell, "Systemic oversimplification limits the potential for human-AI partnership," *IEEE Access*, vol. 9, pp. 70242–70260, 2021.
- [7] Headquarters, Department of the Army, *Army Doctrine Publication 5-0: The Operations Process*, 2019.
- [8] K. McDowell, E. Novoseller, A. Madison, V. Goecks, and C. Kelshaw, "Re-envisioning command and control," in *International Conference on Military Communication and Information Systems*, 2024.
- [9] M. Farmer, "Four-dimensional planning at the speed of relevance: Artificial-intelligence-enabled military decision-making process," *Military Review*, pp. 64–73, Nov-Dec 2022.
- [10] M. Schadd, A. M. Sternheim, R. Blankendaal, M. van der Kaaij, and O. Visker, "How a machine can understand the command intent," *The Journal of Defense Modeling and Simulation*, pp. 1–18, 2022.
- [11] O. Blog, "Introducing ChatGPT," *Internet: https://openai.com/blog/chatgpt*, 2022.
- [12] B. D. Argall, S. Chernova, M. Veloso, and B. Browning, "A survey of robot learning from demonstration," *Robotics and autonomous systems*, vol. 57, no. 5, pp. 469–483, 2009.
- [13] V. G. Goecks, G. M. Gremillion, V. J. Lawhern, J. Valasek, and N. R. Waytowich, "Efficiently combining human demonstrations and interventions for safe training of autonomous systems in real-time," in *AAAI conference on artificial intelligence*, 2019.
- [14] P. Sharma, B. Sundaralingam, V. Blukis, C. Paxton, T. Hermans, A. Torralba, J. Andreas, and D. Fox, "Correcting robot plans with natural language feedback," *arXiv preprint arXiv:2204.05186*, 2022.
- [15] V. G. Goecks and N. R. Waytowich, "DisasterResponseGPT: Large language models for accelerated plan of action development in disaster response scenarios," in *Workshop on Challenges in Deployable Generative AI at International Conference on Machine Learning*, 2023.
- [16] P. F. Christiano, J. Leike, T. Brown, M. Martic, S. Legg, and D. Amodei, "Deep reinforcement learning from human preferences," *Advances in neural information processing systems*, 2017.
- [17] K. Lee, L. M. Smith, and P. Abbeel, "PEBBLE: Feedback-efficient interactive reinforcement learning via relabeling experience and unsupervised pre-training," in *International Conference on Machine Learning*, 2021.
- [18] N. Stiennon, L. Ouyang, J. Wu, D. Ziegler, R. Lowe, C. Voss, A. Radford, D. Amodei, and P. F. Christiano, "Learning to summarize with human feedback," *Advances in Neural Information Processing Systems*, 2020.
- [19] M. Matta, G. C. Cardarilli, L. Di Nunzio, R. Fazzolari, D. Giardino, M. Re, F. Silvestri, and S. Spanò, "Q-RTS: a real-time swarm intelligence based on multi-agent q-learning," *Electronics Letters*, vol. 55, no. 10, pp. 589–591, 2019.
- [20] W. Zhang, X. Wang, J. Shen, and M. Zhou, "Model-based multi-agent policy optimization with adaptive opponent-wise rollouts," in *International Joint Conference on Artificial Intelligence*, 2021.
- [21] R. Wang, J. C. Kew, D. Lee, T.-W. Lee, T. Zhang, B. Ichter, J. Tan, and A. Faust, "Model-based reinforcement learning for decentralized multiagent rendezvous," in *Conference on Robot Learning*, 2021.
- [22] I. Elleuch, A. Pourranjbar, and G. Kaddoum, "A novel distributed multi-agent reinforcement learning algorithm against jamming attacks," *IEEE Communications Letters*, vol. 25, no. 10, pp. 3204–3208, 2021.
- [23] O. Vinyals, I. Babuschkin, W. M. Czarnecki, M. Mathieu, A. Dudzik, J. Chung, D. H. Choi, R. Powell, T. Ewalds, P. Georgiev, et al., "Grandmaster level in StarCraft II using multi-agent reinforcement learning," *Nature*, vol. 575, no. 7782, pp. 350–354, 2019.
- [24] S. A. Wu, R. E. Wang, J. A. Evans, J. B. Tenenbaum, D. C. Parkes, and M. Kleiman-Weiner, "Too many cooks: Bayesian inference for coordinating multi-agent collaboration," *Topics in Cognitive Science*, vol. 13, no. 2, pp. 414–432, 2021.
- [25] A. Bakhtin, N. Brown, E. Dinan, G. Farina, C. Flaherty, D. Fried, A. Goff, J. Gray, H. Hu, et al., "Human-level play in the game of diplomacy by combining language models with strategic reasoning," *Science*, vol. 378, no. 6624, pp. 1067–1074, 2022.
- [26] H. M. Le, Y. Yue, P. Carr, and P. Lucey, "Coordinated multi-agent imitation learning," in *International Conference on Machine Learning*, 2017.
- [27] J. Song, H. Ren, D. Sadigh, and S. Ermon, "Multi-agent generative adversarial imitation learning," *Advances in neural information processing systems*, 2018.
- [28] L. Yu, J. Song, and S. Ermon, "Multi-agent adversarial inverse reinforcement learning," in *International Conference on Machine Learning*, 2019.
- [29] L. Ouyang, J. Wu, X. Jiang, D. Almeida, C. Wainwright, P. Mishkin, C. Zhang, S. Agarwal, K. Slama, A. Ray, et al., "Training language models to follow instructions with human feedback," *Advances in neural information processing systems*, 2022.
- [30] S. Casper, X. Davies, C. Shi, T. Krendl Gilbert, J. Scheurer, J. Rando Ramirez, R. Freedman, T. Korbak, D. Lindner, P. Freire, et al., "Open problems and fundamental limitations of reinforcement learning from human feedback," *Transactions on Machine Learning Research*, 2023.
- [31] F. Parisi, F. Strino, B. Nadler, and Y. Kluger, "Ranking and combining multiple predictors without labeled data," *Proceedings of the National Academy of Sciences*, vol. 111, no. 4, pp. 1253–1258, 2014.
- [32] Y. Sakata, Y. Baba, and H. Kashima, "Crownn: Human-in-the-loop network with crowd-generated inputs," in *IEEE International Conference on Acoustics, Speech and Signal Processing*, 2019.
- [33] Y. Zhang, X. Chen, D. Zhou, and M. I. Jordan, "Spectral methods meet EM: A provably optimal algorithm for crowdsourcing," *Advances in neural information processing systems*, 2014.
- [34] N. B. Shah, S. Balakrishnan, and M. J. Wainwright, "A permutation-based model for crowd labeling: Optimal estimation and robustness," *IEEE Transactions on Information Theory*, vol. 67, no. 6, pp. 4162–4184, 2020.
- [35] G. Zhang and H. Kashima, "Batch reinforcement learning from crowds," in *Joint European Conference on Machine Learning and Knowledge Discovery in Databases*, 2022.
- [36] D. Chhan, E. Novoseller, and V. J. Lawhern, "Crowd-PrefRL: Preference-based reward learning from crowds," *arXiv preprint arXiv:2401.10941*, 2024.
- [37] S. Chakraborty, J. Qiu, H. Yuan, A. Koppel, F. Huang, D. Manocha, A. S. Bedi, and M. Wang, "MaxMin-RLHF: Towards equitable alignment of large language models with diverse human preferences," *arXiv preprint arXiv:2402.08925*, 2024.
- [38] A. Siththaranjan, C. Laidlaw, and D. Hadfield-Menell, "Distributional preference learning: Understanding and accounting for hidden context in RLHF," *arXiv preprint arXiv:2312.08358*, 2023.

- [39] M. Beliaev, A. Shih, S. Ermon, D. Sadigh, and R. Pedarsani, "Imitation learning by estimating expertise of demonstrators," in *International Conference on Machine Learning*, 2022.
- [40] P. Dayan and G. E. Hinton, "Feudal reinforcement learning," *Advances in neural information processing systems*, 1992.
- [41] R. Parr and S. Russell, "Reinforcement learning with hierarchies of machines," *Advances in neural information processing systems*, 1997.
- [42] A. Pashevich, D. Hafner, J. Davidson, R. R. Sukthankar, and C. Schmid, "Modulated policy hierarchies," in *Deep Reinforcement Learning Workshop at Neural Information Processing Systems*, 2018.
- [43] T. G. Dietterich, "Hierarchical reinforcement learning with the MAXQ value function decomposition," *Journal of artificial intelligence research*, vol. 13, pp. 227–303, 2000.
- [44] T. D. Kulkarni, K. Narasimhan, A. Saeedi, and J. Tenenbaum, "Hierarchical deep reinforcement learning: Integrating temporal abstraction and intrinsic motivation," *Advances in neural information processing systems*, 2016.
- [45] A. S. Vezhnevets, S. Osindero, T. Schaul, N. Heess, M. Jaderberg, D. Silver, and K. Kavukcuoglu, "Feudal networks for hierarchical reinforcement learning," in *International Conference on Machine Learning*, 2017.
- [46] S. Sukhbaatar, E. Denton, A. Szlam, and R. Fergus, "Learning goal embeddings via self-play for hierarchical reinforcement learning," *arXiv preprint arXiv:1811.09083*, 2018.
- [47] K. Frans, J. Ho, X. Chen, P. Abbeel, and J. Schulman, "Meta learning shared hierarchies," *arXiv preprint arXiv:1710.09767*, 2017.
- [48] P.-L. Bacon, J. Harb, and D. Precup, "The option-critic architecture," in *AAAI conference on artificial intelligence*, 2017.
- [49] C. Tessler, S. Givony, T. Zahavy, D. Mankowitz, and S. Mannor, "A deep hierarchical approach to lifelong learning in Minecraft," in *AAAI conference on artificial intelligence*, 2017.
- [50] O. Nachum, S. S. Gu, H. Lee, and S. Levine, "Data-efficient hierarchical reinforcement learning," *Advances in neural information processing systems*, 2018.
- [51] A. Brohan, Y. Chebotar, C. Finn, K. Hausman, A. Herzog, D. Ho, J. Ibarz, A. Irpan, E. Jang, R. Julian, *et al.*, "Do as I can, not as I say: Grounding language in robotic affordances," in *Conference on Robot Learning*, 2023.
- [52] D. Driess, F. Xia, M. S. Sajjadi, C. Lynch, A. Chowdhery, B. Ichter, A. Wahid, J. Tompson, Q. Vuong, T. Yu, *et al.*, "PaLM-E: An embodied multimodal language model," in *International Conference on Machine Learning*, 2023.
- [53] G. Wang, Y. Xie, Y. Jiang, A. Mandlkar, C. Xiao, Y. Zhu, L. Fan, and A. Anandkumar, "Voyager: An open-ended embodied agent with large language models," *arXiv preprint arXiv:2305.16291*, 2023.
- [54] C. Gebhardt, A. Oulasvirta, and O. Hilliges, "Hierarchical reinforcement learning as a model of human task interleaving," *Computational Brain and Behavior*, 2021.
- [55] R. Pinsler, R. Akrou, T. Osa, J. Peters, and G. Neumann, "Sample and feedback efficient hierarchical reinforcement learning from human preferences," in *IEEE International Conference on Robotics and Automation*, 2018.
- [56] N. Bougie and R. Ichise, "Hierarchical learning from human preferences and curiosity," *Applied Intelligence*, pp. 1–21, 2022.
- [57] J. Gehring, G. Synnaeve, A. Krause, and N. Usunier, "Hierarchical skills for efficient exploration," *Advances in Neural Information Processing Systems*, 2021.
- [58] X. Wang, K. Lee, K. Hakhamaneshi, P. Abbeel, and M. Laskin, "Skill preferences: Learning to extract and execute robotic skills from human feedback," in *Conference on Robot Learning*, 2022.
- [59] K. J. Friston, M. J. Ramstead, A. B. Kiefer, A. Tschantz, C. L. Buckley, M. Albarracín, R. J. Pitliya, C. Heins, B. Klein, B. Millidge, *et al.*, "Designing ecosystems of intelligence from first principles," *Collective Intelligence*, vol. 3, no. 1, pp. 1–19, 2024.
- [60] P. C. Heredia and S. Mou, "Distributed multi-agent reinforcement learning by actor-critic method," *International Federation of Automatic Control*, vol. 52, no. 20, pp. 363–368, 2019.
- [61] K. Zhang, Z. Yang, and T. Başar, "Decentralized multi-agent reinforcement learning with networked agents: Recent advances," *Frontiers of Information Technology & Electronic Engineering*, vol. 22, no. 6, pp. 802–814, 2021.
- [62] W. Kim, J. Park, and Y. Sung, "Communication in multi-agent reinforcement learning: Intention sharing," in *International Conference on Learning Representations*, 2020.
- [63] K. Kondo, J. Tordesillas, R. Figueroa, J. Rached, J. Merkel, P. C. Lusk, and J. P. How, "Robust MADER: Decentralized and asynchronous multiagent trajectory planner robust to communication delay," in *IEEE International Conference on Robotics and Automation*, 2023.
- [64] H.-T. Wai, Z. Yang, Z. Wang, and M. Hong, "Multi-agent reinforcement learning via double averaging primal-dual optimization," *Advances in Neural Information Processing Systems*, 2018.
- [65] R. Raileanu, E. Denton, A. Szlam, and R. Fergus, "Modeling others using oneself in multi-agent reinforcement learning," in *International conference on machine learning*, 2018.
- [66] S. V. Albrecht and P. Stone, "Autonomous agents modelling other agents: A comprehensive survey and open problems," *Artificial Intelligence*, vol. 258, pp. 66–95, 2018.
- [67] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, Ł. Kaiser, and I. Polosukhin, "Attention is all you need," *Advances in neural information processing systems*, 2017.
- [68] J. D. Kenton, M.-W. Chang, and L. K. Toutanova, "BERT: Pre-training of deep bidirectional transformers for language understanding," in *Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, 2019.
- [69] C. Raffel, N. Shazeer, A. Roberts, K. Lee, S. Narang, M. Matena, Y. Zhou, W. Li, and P. J. Liu, "Exploring the limits of transfer learning with a unified text-to-text transformer," *The Journal of Machine Learning Research*, vol. 21, no. 1, pp. 5485–5551, 2020.
- [70] T. Brown, B. Mann, N. Ryder, M. Subbiah, J. D. Kaplan, P. Dhariwal, A. Neelakantan, P. Shyam, G. Sastry, A. Askell, *et al.*, "Language models are few-shot learners," *Advances in neural information processing systems*, 2020.
- [71] T. Le Scao and A. M. Rush, "How many data points is a prompt worth?," in *Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, 2021.
- [72] S. Min, X. Lyu, A. Holtzman, M. Artetxe, M. Lewis, H. Hajishirzi, and L. Zettlemoyer, "Rethinking the role of demonstrations: What makes in-context learning work?," *arXiv preprint arXiv:2202.12837*, 2022.
- [73] S. M. Xie, A. Raghunathan, P. Liang, and T. Ma, "An explanation of in-context learning as implicit Bayesian inference," in *International Conference on Learning Representations*, 2021.
- [74] J. Luketina, N. Nardelli, G. Farquhar, J. Foerster, J. Andreas, E. Grefenstette, S. Whiteson, and T. Rocktäschel, "A survey of reinforcement learning informed by natural language," in *International Joint Conference on Artificial Intelligence*, 2019.
- [75] Y. Du, O. Watkins, Z. Wang, C. Colas, T. Darrell, P. Abbeel, A. Gupta, and J. Andreas, "Guiding pretraining in reinforcement learning with large language models," in *International Conference on Machine Learning*, 2023.
- [76] S. D. Ramchurn, S. Stein, and N. R. Jennings, "Trustworthy human-AI partnerships," *Isience*, vol. 24, no. 8, 2021.
- [77] M. P. Schadd, T. A. Schoonderwoerd, K. van den Bosch, O. H. Visker, and T. Haije, "'I'm afraid I can't do that, Dave': getting to know your buddies in a human-agent team," *Systems*, vol. 10, no. 1, p. 15, 2022.
- [78] J. Van Diggelen, M. Neerincx, M. Peeters, and J. M. Schraagen, "Developing effective and resilient human-agent teamwork using team design patterns," *IEEE intelligent systems*, vol. 34, no. 2, pp. 15–24, 2018.
- [79] J. van Diggelen and M. Johnson, "Team design patterns," in *International Conference on Human-Agent Interaction*, 2019.
- [80] A. Schulte, D. Donath, and D. S. Lange, "Design patterns for human-cognitive agent teaming," in *Engineering Psychology and Cognitive Ergonomics: International Conference, Held as Part of HCI International*, 2016.
- [81] J. van der Waa, S. Verdult, K. van den Bosch, J. van Diggelen, T. Haije, B. van der Stigchel, and I. Cocu, "Moral decision making in human-agent teams: Human control and the role of explanations," *Frontiers in Robotics and AI*, vol. 8, pp. 1–20, 2021.
- [82] L. Onnasch and E. Roesler, "A taxonomy to structure and analyze human-robot interaction," *International Journal of Social Robotics*, vol. 13, no. 4, pp. 833–849, 2021.
- [83] D. Siemon, "Elaborating team roles for artificial intelligence-based teammates in human-AI collaboration," *Group Decision and Negotiation*, vol. 31, no. 5, pp. 871–912, 2022.
- [84] S. D. Ramchurn, T. D. Huynh, F. Wu, Y. Ikuno, J. Flann, L. Moreau, J. E. Fischer, W. Jiang, T. Rodden, E. Simpson, *et al.*, "A disaster re-

- sponse system based on human-agent collectives,” *Journal of Artificial Intelligence Research*, vol. 57, pp. 661–708, 2016.
- [85] K. A. Hoff and M. Bashir, “Trust in automation: Integrating empirical evidence on factors that influence trust,” *Human factors*, vol. 57, no. 3, pp. 407–434, 2015.
- [86] C. D. Wickens, B. A. Clegg, A. Z. Vieane, and A. L. Sebok, “Complacency and automation bias in the use of imperfect automation,” *Human factors*, vol. 57, no. 5, pp. 728–739, 2015.
- [87] J. D. Lee and K. A. See, “Trust in automation: Designing for appropriate reliance,” *Human factors*, vol. 46, no. 1, pp. 50–80, 2004.
- [88] C. J. Johnson, M. Demir, N. J. McNeese, J. C. Gorman, A. T. Wolff, and N. J. Cooke, “The impact of training on human–autonomy team communications and trust calibration,” *Human factors*, vol. 65, no. 7, pp. 1554–1570, 2023.
- [89] R. W. Andrews, J. M. Lilly, D. Srivastava, and K. M. Feigh, “The role of shared mental models in human-AI teams: a theoretical review,” *Theoretical Issues in Ergonomics Science*, vol. 24, no. 2, pp. 129–175, 2023.
- [90] J. D. Weisz, M. Muller, S. Houde, J. Richards, S. I. Ross, F. Martinez, M. Agarwal, and K. Talamadupula, “Perfection not required? Human-AI partnerships in code translation,” in *International Conference on Intelligent User Interfaces*, 2021.
- [91] S. A. Kusumastuti, K. A. Pollard, A. H. Oiknine, B. Dalangin, T. R. Raber, and B. T. Files, “Practice improves performance of a 2D uncertainty integration task within and across visualizations,” *IEEE Transactions on Visualization and Computer Graphics*, 2022.
- [92] M. Kay, T. Kola, J. R. Hullman, and S. A. Munson, “When (ish) is my bus? User-centered visualizations of uncertainty in everyday, mobile predictive systems,” in *CHI conference on human factors in computing systems*, 2016.
- [93] N. J. McNeese, B. G. Schelble, L. B. Canonico, and M. Demir, “Who/what is my teammate? Team composition considerations in human-AI teaming,” *IEEE Transactions on Human-Machine Systems*, vol. 51, no. 4, pp. 288–299, 2021.
- [94] S. Campbell, R. Nguyen, E. Bonsignore, B. Carter, and C. Neubauer, “Defining and modeling AI technical fluency for effective human machine interaction,” in *Human Factors in Robots, Drones and Unmanned Systems*, AHFE International, 2023.
- [95] J. Duncan, D. Chylinski, D. J. Mitchell, and A. Bhandari, “Complexity and compositionality in fluid intelligence,” *Proceedings of the National Academy of Science*, vol. 114, no. 20, pp. 5295–5299, 2017.
- [96] K. A. Pollard, B. T. Files, A. H. Oiknine, and B. Dalangin, “How to prepare for rapidly evolving technology: Focus on adaptability,” tech. rep., ARL-TR-9432. US Combat Capabilities Development Command, Army Research Laboratory, Aberdeen Proving Ground, United States, 2022.
- [97] S. Milani, N. Topin, M. Veloso, and F. Fang, “A survey of explainable reinforcement learning,” *arXiv preprint arXiv:2202.08434*, 2022.
- [98] R. Confalonieri, L. Coba, B. Wagner, and T. R. Besold, “A historical perspective of explainable artificial intelligence,” *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, vol. 11, no. 1, p. e1391, 2021.
- [99] J. Cha, M. Barnes, and J. Y. Chen, “Visualization techniques for transparent human-agent interface designs,” tech. rep., ARL-TR-8674. US Combat Capabilities Development Command Army Research Laboratory Aberdeen Proving Ground, United States, 2019.
- [100] Y. Zhang, Q. V. Liao, and R. K. Bellamy, “Effect of confidence and explanation on accuracy and trust calibration in AI-assisted decision making,” in *conference on fairness, accountability, and transparency*, 2020.
- [101] G. Bansal, T. Wu, J. Zhou, R. Fok, B. Nushi, E. Kamar, M. T. Ribeiro, and D. Weld, “Does the whole exceed its parts? The effect of AI explanations on complementary team performance,” in *CHI Conference on Human Factors in Computing Systems*, 2021.
- [102] H. Vasconcelos, M. Jörke, M. Grunde-McLaughlin, T. Gerstenberg, M. S. Bernstein, and R. Krishna, “Explanations can reduce overreliance on AI systems during decision-making,” *Proceedings of the ACM on Human-Computer Interaction*, vol. 7, no. CSCW1, pp. 1–38, 2023.
- [103] T. Miller, “Explanation in artificial intelligence: Insights from the social sciences,” *Artificial intelligence*, vol. 267, pp. 1–38, 2019.
- [104] M. A. Neerinx, J. van der Waa, F. Kaptein, and J. van Diggelen, “Using perceptual and cognitive explanations for enhanced human-agent team performance,” in *Engineering Psychology and Cognitive Ergonomics: International Conference, Held as Part of HCI International*, 2018.
- [105] S. Reddy, A. Dragan, S. Levine, S. Legg, and J. Leike, “Learning human objectives by evaluating hypothetical behavior,” in *International Conference on Machine Learning*, 2020.
- [106] Y. Liu, G. Datta, E. Novoseller, and D. S. Brown, “Efficient preference-based reinforcement learning using learned dynamics models,” in *IEEE International Conference on Robotics and Automation*, 2023.
- [107] L. Padilla, S. C. Castro, and H. Hosseinpour, “A review of uncertainty visualization errors: Working memory as an explanatory theory,” *Psychology of Learning and Motivation*, vol. 74, pp. 275–315, 2021.
- [108] A. Kale, F. Nguyen, M. Kay, and J. Hullman, “Hypothetical outcome plots help untrained observers judge trends in ambiguous data,” *IEEE transactions on visualization and computer graphics*, vol. 25, no. 1, pp. 892–902, 2018.
- [109] Z. Ghahramani, “Probabilistic machine learning and artificial intelligence,” *Nature*, vol. 521, no. 7553, pp. 452–459, 2015.
- [110] Y. Gal and Z. Ghahramani, “Dropout as a Bayesian approximation: Representing model uncertainty in deep learning,” in *International Conference on Machine Learning*, 2016.
- [111] S. Bubeck, V. Chandrasekaran, R. Eldan, J. Gehrke, E. Horvitz, E. Kamar, P. Lee, Y. T. Lee, Y. Li, S. Lundberg, et al., “Sparks of artificial general intelligence: Early experiments with GPT-4,” *arXiv preprint arXiv:2303.12712*, 2023.
- [112] G. Marcus, E. Davis, and S. Aaronson, “A very preliminary analysis of DALL-E 2,” *arXiv preprint arXiv:2204.13807*, 2022.
- [113] M. E. Taylor, N. Carboni, A. Fachantidis, I. Vlahavas, and L. Torrey, “Reinforcement learning agents providing advice in complex video games,” *Connection Science*, vol. 26, no. 1, pp. 45–63, 2014.
- [114] M. Samadi, P. Talukdar, M. Veloso, and T. Mitchell, “AskWorld: Budget-sensitive query evaluation for knowledge-on-demand,” in *International Joint Conference on Artificial Intelligence*, 2015.
- [115] E. Ilhan, J. Gow, and D. Perez-Liebana, “Teaching on a budget in multi-agent deep reinforcement learning,” in *IEEE Conference on Games*, 2019.
- [116] Z. Xu, K. Q. Weinberger, and O. Chapelle, “The greedy miser: Learning under test-time budgets,” in *International Conference on Machine Learning*, 2012.
- [117] S. Basu, S. Gutstein, B. Lance, and S. Shakkottai, “Pareto optimal streaming unsupervised classification,” in *International Conference on Machine Learning*, 2019.
- [118] J. Song, Y. Chen, and Y. Yue, “A general framework for multi-fidelity Bayesian optimization with Gaussian processes,” in *International Conference on Artificial Intelligence and Statistics*, 2019.
- [119] D. Reifsteck, T. Engesser, R. Mattmüller, and B. Nebel, “Epistemic multi-agent planning using Monte-Carlo tree search,” in *Advances in Artificial Intelligence: German Conference on AI*, 2019.
- [120] O. Kaduri, E. Boyarski, and R. Stern, “Algorithm selection for optimal multi-agent pathfinding,” in *International Conference on Automated Planning and Scheduling*, 2020.
- [121] K. Kasaura, M. Nishimura, and R. Yonetani, “Prioritized safe interval path planning for multi-agent pathfinding with continuous time on 2D roadmaps,” *IEEE Robotics and Automation Letters*, vol. 7, no. 4, pp. 10494–10501, 2022.
- [122] A. Wong, T. Bäck, A. V. Kononova, and A. Plaat, “Deep multiagent reinforcement learning: Challenges and directions,” *Artificial Intelligence Review*, vol. 56, no. 6, pp. 5023–5056, 2023.
- [123] J. Derleth, “Enhancing interoperability: The foundation for effective NATO operations,” *NATO Review*, 2015.
- [124] P. Mangalraj, V. Sivakumar, S. Karthick, V. Haribaabu, S. Ramraj, and D. J. Samuel, “A review of multi-resolution analysis (MRA) and multi-geometric analysis (MGA) tools used in the fusion of remote sensing images,” *Circuits, Systems, and Signal Processing*, vol. 39, pp. 3145–3172, 2020.
- [125] K. Zheng, Y. Sung, G. Konidaris, and S. Tellex, “Multi-resolution POMDP planning for multi-object search in 3D,” in *IEEE/RSJ International Conference on Intelligent Robots and Systems*, 2021.
- [126] W. Wang, Y. Liu, R. Srikant, and L. Ying, “3M-RL: multi-resolution, multi-agent, mean-field reinforcement learning for autonomous UAV routing,” *IEEE Transactions on Intelligent Transportation Systems*, vol. 23, no. 7, pp. 8985–8996, 2021.