



Broad Agency Announcement Information Processing Techniques Office

Decentralized Artificial Intelligence through Controlled
Emergence (DICE)

HR001126S0010

June 10, 2026

Amendment 01

This publication constitutes a Broad Agency Announcement (BAA) as contemplated in Federal Acquisition Regulation (FAR) 6.102(d)(2) and 35.016 and 2 CFR § 200.203. Any resultant award negotiations will follow all pertinent law and regulation, and any negotiations and/or awards for procurement contracts will use procedures under FAR 15.4, Contract Pricing, as specified in the BAA.

Amendment 01

Amendment 01 is being issued to update language regarding the Cybersecurity Maturity Model Certification Requirements. All changes are highlighted.

TABLE OF CONTENTS

TABLE OF CONTENTS	2
1. Overview Information	3
2. Funding Opportunity Description.....	5
2.1 Program Goal.....	5
2.2 Program Background.....	6
2.3 Program Introduction.....	8
2.4 Program Structure	10
2.5 Program Metrics	20
2.6 Program Schedule, Meetings, and Milestones.....	22
2.6.2 Program Meetings.....	23
2.7 Anticipated Deliverables	23
2.8 Intellectual Property and Data Rights.....	24
3. Security	25
3.1 CUI	25
3.2 DARPA Fundamental Research Risk-Based Security Review (FRRBS) Background.....	25
3.2.1 DARPA FRRBS Required Documents	25
3.3 Cybersecurity Maturity Model Certification (CMMC) Requirements	26
4. Submission Information.....	26
5. AbstractS (Strongly Encouraged).....	27
5.1 Abstract Submission Requirements.....	27
5.2 Abstract Content and Formatting Requirements	27
5.3 Abstract Review.....	27
5.4 Abstract Feedback	28
6. ProposalS	28
7. Proposal Evaluation Criteria.....	29
8. Special Considerations.....	30
9. References.....	32

1. OVERVIEW INFORMATION

- **Federal Agency Name:** Defense Advanced Research Projects Agency (DARPA), Information Processing Techniques Office (IPTO)
- **Funding Opportunity Title:** Decentralized Artificial Intelligence through Controlled Emergence (DICE)
- **Funding Opportunity Number:** HR001126S0010
- **Announcement Type:** Initial Announcement
- **Assistance Listing Number:** 12.910
- **Dates/Time:** *All Times are Eastern Time Zone (ET)*
 - **Proposers Day:** May 29, 2026
 - **Posting Date:** June 10, 2026
 - **Proposal Abstract Due Date:** June 30, 2026 at 2:00 PM
 - **Question Submittal Closed:** August 18, 2026 at 2:00 PM
 - **Proposal Due Date:** August 25, 2026 at 2:00 PM
- **Anticipated Individual Awards:** Multiple awards are anticipated.
- **Types of Instruments that may be Awarded:** Procurement Contract, Other Transaction (OT) for Prototype Agreement (10 U.S.C. § 4022), OT for Research Agreement (10 U.S.C. § 4021), Cooperative Agreement
- **Resource Sharing Requirements:** In accordance with 10 U.S.C. § 4021, and 10 U.S.C. § 4022, resource sharing may be required for OT for Prototype or OT for Research Agreements
- **NAICS Code:** 541715
- **Questions and Answers (Q&As):** All administrative, technical, and award questions should be emailed to the Agency Point of Contact (POC) stated below. All questions must be written in English and must include the name of a designated point of contact. Q&A documents will be published on the DICE program page at:
<https://www.darpa.mil/research/programs/decentralized-artificial-intelligence-through-controlled-emergence>
- **Agency POC:** The BAA Coordinator for this effort may be reached at: DICE@darpa.mil
 DARPA / IPTO
 ATTN: HR001126S0010
 675 North Randolph Street
 Arlington, VA 22203-2114
- **Attachments:**
 - DICE Controlled Unclassified Information (CUI) Guide
 - A1 – DICE Abstract Template
 - P1 – DICE Volume I Proposal Template (*Technical and Management*)
 - P2 – DICE Volume II Proposal Template (*Cost*)
 - P3 – DICE Proposal Summary Slide Template
 - P4 - DARPA Standard Cost Proposal Spreadsheet
 - CF Biographical Sketch
 - CF Biographical Sketch Sample

- CF Other Support
- CF Other Support Sample
- Associate Contractor Agreement (*provided for informational purposes only*)

2. FUNDING OPPORTUNITY DESCRIPTION

The Defense Advanced Research Projects Agency (DARPA) is soliciting innovative proposals in the following technical areas: artificial intelligence (AI), control theory, formal methods, and game theory to develop decentralized multi-agentic AI collectives. Proposed research should investigate innovative approaches that enable revolutionary advances in science, devices, or systems. Specifically excluded is research that primarily results in evolutionary improvements to the existing state of practice.

2.1 PROGRAM GOAL

Decentralized Artificial Intelligence through Controlled Emergence (DICE) aims to develop the theory and algorithms for decentralized coordination and local inference control to enable a scalable, adaptive, and resilient collective of heterogeneous AI agents that can autonomously execute sustained, long-time-horizon missions in contested environments while remaining under control. The goal is to move beyond brittle, centralized orchestration through workflows and ad hoc agent compositions of AI agents by harnessing self-organization while constraining emergent behavior to remain predictable, doctrinally consistent, mission-aligned in adherence with user intent, and resilient to benign failures and adversarial compromise.

This fundamental shift to decentralized intelligence envisioned by DICE is motivated by the need to transform the existing approach to solving complex problems that require multiskilled teams operating in dynamic environments, such as responding to natural disasters, advancing scientific discovery, or engineering complex systems. Future conflicts will unfold at machine speed, and the current centralized planning and scripted orchestration approach is too slow, too rigid, and too predictable for the hyper-dynamic, contested environments of the future. Surprise, adaptability, and resilience will depend on the ability to compose and recompose AI capabilities on demand. DICE is built for such autonomy-powered, high-tempo battlespace. Imagine a scalable, adaptable, and resilient force of autonomous systems that can execute complex, long-term missions in contested environments without a human explicitly directing every action while remaining under human control and in adherence with user's intent. These autonomous forces will be able to sense a change in the battlespace and adapt accordingly. It is a future where success is measured by the relentless, creative ability of an autonomous collective to execute missions in the face of chaos and uncertainty. This is a profound and necessary departure from the past.

To enable this vision, DICE seeks to build a highly scalable decentralized coordination architecture that decomposes complex goals and fuses distributed situational context across specialized agents over long-time-horizons. The architecture will support multistep interaction among agents, iterative information gathering under partial observability, adaptive strategy refinement as the environment evolves, and resilience to broken links, conflicting information, agents going rogue by developing instrumental goals misaligned with mission, and the loss, failure, or adversarial compromise of individual agents. Through peer-to-peer coordination and resilient distributed consensus, agents will be able to self-organize into teams, allocate roles, and fuse conflicting, incomplete, or adversarial information without waiting for a central human or machine planner. In DICE, AI capabilities will emerge from the interactions among heterogeneous agents, extending beyond the intelligence of individual agents.

At the same time, DICE recognizes that decentralized coordination alone is insufficient. Sustained long-horizon missions require local control of the individual agent as well as governance of the collective. DICE therefore seeks to develop local inference-time control methods that maintain role coherence across multi-turn interactions and constrain collective-level emergent behavior, while still preserving the cognitive agility needed to discover new courses of action in a changing environment. In DICE, intelligence is allowed to emerge from the interactions of heterogeneous agents, but that emergence is bounded by architectural controls that preserve doctrine, suppress agent misbehavior, maintain mission focus, and sustain long-horizon coordination.

DICE aims to enable this future by creating a local adaptor that can interface with any agent and is responsible for peer-to-peer coordination and local control. The agents in DICE can be heterogeneous and use different foundation models. This approach mirrors the principles of decentralized self-organization¹ that underpin the internet's own scalability and resilience, where robust global behavior emerges from well-designed simple, local rules. DICE aims to create an agile dynamic internet of agents at the cognitive level of interaction. Self-organization using local adaptors at the cognitive level requires distributed mission decomposition into tasks for each agent, robust fusion of information from agents that have partial and noisy observability, the absence of any synchronized shared world model across all agents, and the consequential need for local decision-making. The program seeks to demonstrate DICE architectures in DoW-relevant use-cases in simulation environments, targeting measurable gains in scalability, adaptability, resilience, and long-horizon role coherence, with the ultimate aim of achieving Observe–Orient–Decide–Act (OODA) loop superiority in future contested operations.

2.2 PROGRAM BACKGROUND

Recent advances in foundation models [1] have accelerated the development of autonomous AI agents capable of reasoning, planning, following instructions, and using external tools across digital and physical domains. Large Language Models (LLMs) [2], Vision-Language Models (VLMs) [3], and Vision-Language-Action Models (VLAs) [4] have enabled a growing ecosystem of single-agent and multi-agent systems for tasks ranging from software engineering [5] and legal support [6] to robotics [7] and cyber operations [8]. These advances, together with progress toward standardizing tool calling (e.g., model context protocol [9]) and agent-to-agent communication (e.g., Agent2Agent protocol [10], Agent Network Protocol [11]), have made AI agents increasingly relevant to commercial, scientific, and military applications. As these systems have matured, developers have increasingly turned from single-agent systems [12,13] to multi-agent systems (MAS) [14-16] to address tasks requiring varied expertise, persistent state, and coordination among specialized roles. In many cases, such specialized teams can outperform a single, larger model on complex tasks by distributing work across agents with distinct functions such as planning, analysis, sensing, testing, or execution. Yet most current MAS are assembled in an ad hoc manner and remain governed by a central planner, router, or manually scripted workflow. This makes them useful in narrow settings with structured tasks, but ill-suited for the uncertainty, scale, and adversarial conditions of contested operations.

The dominant weakness of today's MAS is central orchestration. This can be a manual workflow implemented in orchestration frameworks such as LangGraph [17], CrewAI [18] or AutoGen [19],

¹ <https://www.nature.com/articles/s44260-025-00031-5>

or the use of another AI orchestrator agent that assigns tasks, and aggregates outputs. This creates a catastrophic single point of failure and a major bottleneck for scale. It constrains the number of agents, roles, and interactions (messages exchanges between the agents) the collective can support. Even when the orchestrator is another AI agent, the mission's growing complexity eventually exceeds the context and inference limits of the underlying foundation model. Further, centralized architecture may be acceptable for predictable enterprise automation, but they are fundamentally mismatched to contested environments where communications can be degraded, information is often incomplete, and rapid decentralized adaptation is essential [20]. Central orchestration also creates a single point of failure and makes the overall MAS less resilient to benign failures and adversarial attacks. These weaknesses of existing MAS architectures manifest in four critical, demonstrable limitations: lack of scalability, limited adaptability, difficulty operating with partial observations, and fragility to adversarial attacks [21-24].

These limitations are not merely engineering challenges that can be addressed through incremental evolutionary improvements or just increase in model size [25, 26] but have more fundamental root causes requiring foundational shift in MAS architecture. Transformer attention scales poorly with long contexts, making it difficult to preserve relevant information over long time horizons or coordinate many specialized agents through a single reasoning bottleneck. As mission complexity grows, centralized coordination becomes slower and more brittle, limiting scalability. When environments change, all relevant information needs to propagate before a revised plan can be formed and distributed, delaying response and reducing the collective's ability to exploit fleeting opportunities or rapidly adapt to external changes or internal disruptions. State-of-the-art (SOTA) MAS also remain weak under partial observability. Because generative models are trained with autoregressive objectives, they often hallucinate or confabulate when critical information is missing rather than explicitly deferring, requesting additional evidence, or waiting for inputs from other agents. This makes them poorly suited for contested missions in which the correct action depends on distributed information gathering and reasoning over ambiguous, conflicting, or incomplete evidence. The fragility of MAS stems from the vulnerability of foundation models to adversarial manipulation, and in centralized architecture, a compromised orchestrator or an infected agent can propagate errors throughout the collective. This fragility is especially problematic for DoW-relevant missions, where benign failures, degraded communications, deceptive inputs, and strategically malicious behavior are expected rather than exceptional. A system that cannot isolate faults, contain compromise, and continue operating after partial loss is not suitable for sustained autonomous use in contested environments. A further limitation is that existing MAS cannot reliably execute long-time-horizon missions. Even when short-horizon performance is strong, individual agents often lose coherence over repeated multi-turn interactions, leading to significant declines in both performance and reliability. This degradation makes current systems poorly suited for missions that require thousands of reasoning steps, repeated coordination among specialized agents, and persistent adherence to role, doctrine, and commander's intent.

At the same time, recent technical progress suggests a path beyond these limitations. Advances in self-organizing systems [27-30] and distributed consensus [31-33] provide a basis for peer-to-peer coordination without centralized control. Conway's Game of Life demonstrates how simple interactions yield emergent complexity; however, a formal science for engineering these rules in complex cognitive units, such as AI agents, is required to achieve the controlled emergence of aligned, self-evolving collective intelligence. Progress in Byzantine-resilient consensus [32-33] and mechanism design [34-35] suggests ways to sustain cooperation and isolate compromised

agents. In parallel, breakthroughs in mechanistic interpretability [36-39], activation steering [40-43], and management of agentic memory and context [44-48] suggest that the internal reasoning of foundation models can be monitored and shaped at inference time without full retraining. Together, these developments make it timely to move beyond brittle orchestration toward a principled architecture for decentralized, controllable collective intelligence.

DICE is motivated by these gaps and opportunities. The program begins with the premise that robust multi-agent AI for contested environments requires solving two coupled problems simultaneously: decentralized coordination of the collective, and local control of the individual agent. DICE will address the first through decentralized self-organization, distributed task decomposition, and resilient consensus, and the second through local inference-time control that preserves role coherence and mission alignment over long horizons. In doing so, DICE seeks to replace brittle centralized orchestration with a new architecture for controlled emergence: a decentralized intelligence collective in which capabilities emerge from the interactions among heterogeneous agents, but that emergence remains bounded by architectural controls that keep the system resilient, predictable, doctrinally consistent, and on mission.

2.3 PROGRAM INTRODUCTION

DICE will aim to develop a decentralized AI architecture suitable for rapidly evolving, unpredictable, and contested environments, that is:

- Scalable: able to support long-horizon complex missions requiring many specialized agents and many turns of interaction.
- Adaptable: able to make rapid decentralized decisions under communication constraints and partial observability of the agents in a rapidly changing environment.
- Resilient: able to tolerate ambiguous or conflicting information, benign dropouts, and adversarial compromise of agents, and resilient against agents going rogue by developing instrumental goals misaligned to the overall mission.

The collective should be able to incorporate heterogeneous, multi-vendor AI agents, such as agents built using different foundation models or different architectures.

As demonstrated by ClawBots [13] and Moltbook [16], a self-evolving AI collective experiences increase in entropy, exhibiting undesirable emergent behavior and reducing cohesion and functionality even in absence of adversarial perturbations. Self-organization in nature overcomes the natural increase in entropy for physical systems and leads to emergence of order. Inspired by this, DICE will seek to expand and leverage the theory of self-organizing systems and distributed consensus algorithms together with recent breakthroughs in controlling the internal reasoning and inference of AI foundation models at inference-time. In DICE, AI agents will aim to use peer-to-peer coordination to self-organize themselves into teams for executing complex missions. This coordination should be robust to failure or compromise of individual agents and use distributed consensus for resilient fusion of context in the presence of conflicting information. The local inference control on each AI agent will ensure role coherence of the individual agents and constrain the emergent behavior of the collective to adhere to user’s intent over the long term and across multiple inference steps. The controller will prevent role-drift ensuring sustained coordination and mission alignment while simultaneously preserving cognitive agility of discovering new courses of actions in a dynamic environment. The coordination and control mechanisms will be

implemented as an adaptor that interfaces locally with each agent as illustrated in Figure 1. The program will demonstrate DICE in DoW-relevant use-cases in simulation environments.

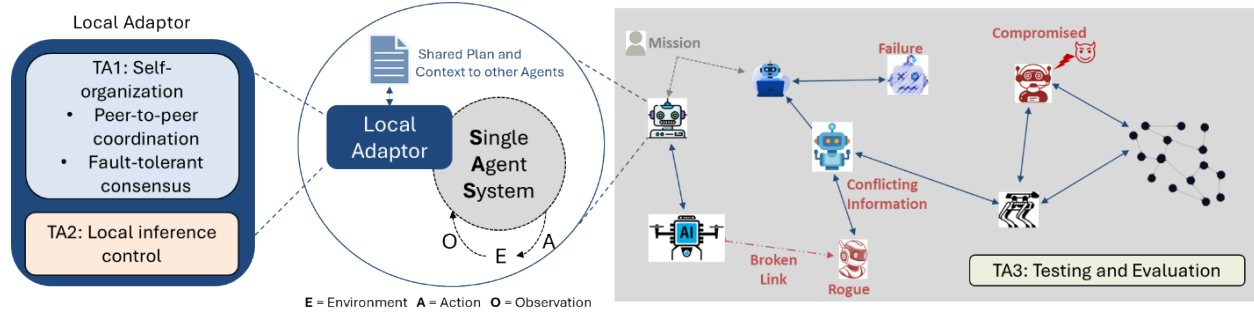


Figure 1. A local adaptor for peer-to-peer coordination and local inference control on each agent in DICE enables a controlled emergence of a scalable, adaptable, and resilient self-organizing collective of AI agents. The DICE adaptor should be easily integrable with heterogeneous AI agents. TA1-TA2 are responsible for developing the adaptor, and TA3 provides a simulation environment for testing with the AI agents and hooks to connect the TA1-TA2 adaptor. All proposals are invited to suggest ideal interface between the adaptor and the single agents. This interface will have different modes and attributes to support agents with open-weight models and those with closed models. For example, the interface to agents with closed models could provide access to memory and other elements of context engineering and the interface to agents with open-weight models could also provide access to internal activations. DICE will converge to a standard adaptor interface early in the program.

The self-organization through peer-to-peer coordination and consensus will be addressed by Technical Area (TA) 1; the control of AI agents will be addressed by TA2; and TA3 will be responsible for testing and evaluating DoW-relevant use-cases. The goals of each of the TAs are described below.

- TA1 teams will pursue different technical approaches for peer-to-peer coordination that include (but are not limited to) distributed auction for decentralized task planning, multi-agent reinforcement learning, and fault-tolerant consensus methods. TA1 coordination and consensus methods aim at improving scalability measured in terms of mission complexity (specified using the number of agents, roles, and the number of interactions/messages between them to execute the mission) and adaptability that is primarily measured in terms of time-to-recovery.
- TA2 teams will pursue diverse strategies for local control that include (but are not limited to) activation steering, agent memory editing, and context engineering. TA2 control methods aim at balancing alignment to mission and roles of each agent while preserving their cognitive agility. This balance can be measured by the diversity of the courses of actions and their sustained mission alignment.
- TA3 team(s) will be responsible for Test and Evaluation (T&E) and organizing evaluation events. TA3 will also be responsible for the development of the simulation environment to be used in T&E including the AI agents to which TA1-TA2 adaptors will connect, the development of use-cases for demonstrations throughout the program, and the evaluation of integrated TA1-TA2 solutions.

There is a close interdependence between coordination (TA1) and control (TA2). Scalable coordination of a large number of agents over longer-time-horizons requires controllable agents that do not lose role coherence over multiple interaction steps. Balancing control and cognitive agility requires the controller to enforce constraints needed for coordination while preserving the capability of discovering new courses of actions by the agents and the collective. The TA1 coordination mechanism will use invariants provided by the TA2 controller, e.g., the distributed

planning horizon in TA1 will depend on the TA2 guarantee on the length of role coherence. In turn, the TA2 control target will be determined by the TA1 coordination requirements.

From an engineering perspective, the TA1-TA2 team will jointly develop the local adaptor that interfaces with individual agents locally to provide the coordination and control needed for self-organization and controlled emergence of scalable, adaptive, and resilient AI collectives. This adaptor will interface with the heterogeneous AI agents in the TA3 (T&E) testing simulation environment. The interface between the agents and the adaptors will be different for agents with open-weight models and agents with closed models. Proposals for all TAs are invited to propose an effective interface most suited for their technical approach. DICE will eventually adopt a standardized interface between the TA1-TA2 adaptor and the TA3 agents that achieves the goal of decentralized coordination and control while maintaining generality and applicability across heterogeneous AI agents, and all performers will use the same interface. TA1-TA2 are encouraged to clearly identify assumptions such as the use of foundation models in TA3 agents or other reasoning capability in the agents when describing their technical approach and defining an interface between the adaptor and the agents.

2.4 PROGRAM STRUCTURE

DICE is a 36-month, 3-phase effort with three TAs. The first phase, spanning nine months, is focused on decentralization and early demonstration of improvement against SOTA orchestrated MAS. The second phase, which consists of 15 months, is focused on adversarial robustness, and the third phase spanning 12 months is focused on scalability. Phase 1 evaluation will compare decentralized DICE collective with SOTA centrally orchestrated MAS. The evaluation in Phases 2 and 3 will use team vs team competition.

Several technical approaches are discussed below to serve as illustrative examples for contextualizing the key challenges within the three TAs. They are not intended to be prescriptive or exhaustive. Competitive proposals will go significantly beyond these illustrative concepts, proposing disruptive, novel techniques or demonstrating groundbreaking refinements to existing state-of-the-art approaches.

Proposals must comprehensively address all aspects of technical area needed as detailed below to attain the program metrics described in Section 2.5 (Program Metrics).

The proposal should not address all three TAs. TA1-TA2 proposals must address TA1 and TA2 together. Proposals for TA3 must address only this TA. Organizations selected for TA3 will not be eligible to perform in TA1 or TA2.

TA1. Self-organization via Peer-to-Peer Coordination and Distributed Consensus

TA1 will aim to develop the theory, algorithms, and implementations required for heterogeneous AI agents to self-organize into mission-effective teams using peer-to-peer coordination, distributed task planning, resilient context fusion, and consensus under benign and adversarial failures. TA1 aims to improve scalability, adaptability, and resilience. Scalability can be measured by mission complexity - number of agents, number of roles and number of agent interactions (messages between the agents) needed to accomplish the mission. Adaptability is measured primarily by time-to-recovery after perturbations, failures, or compromises. Resilience can be measured by mission success and coordination stability under benign failures, Byzantine failures, deception, and conflicting information.

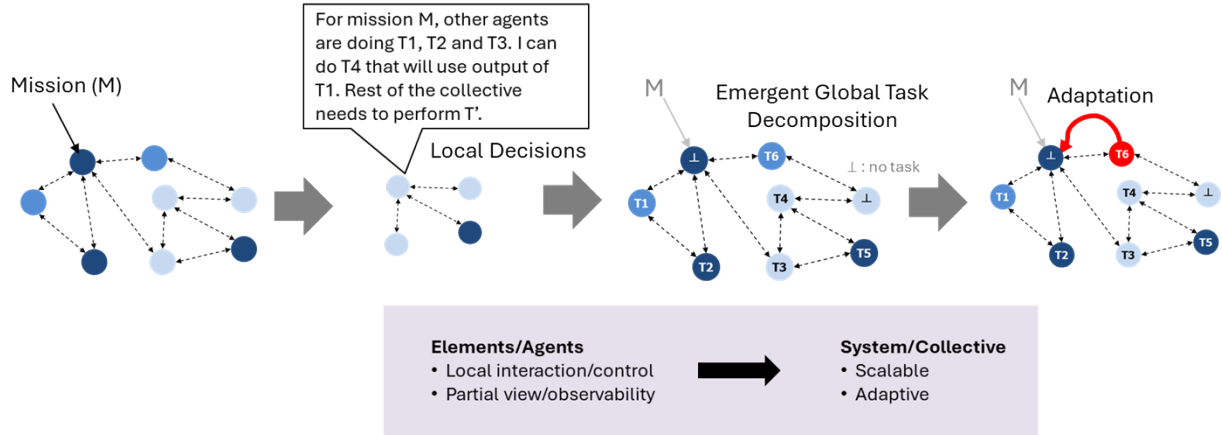


Figure 2. An example of a peer-to-peer coordination protocol for distributed decomposition of missions into tasks for each agent can be based on auction where each agent can bid to execute tasks that are relevant to the mission and within the skills of the agent while broadcasting the remaining mission to the collective for other agents to bid. Based on these local bids and partial observability of the agents, the collective can reach consensus on the ideal task decomposition without a central orchestrator. Failure of an agent can lead to local auction of the failed agent's task without full replanning for the entire collective.

TA1 performers will pursue decentralized coordination and consensus mechanisms, including but not limited to distributed auction and market-based task allocation, consensus-based fusion algorithms, multi-agent reinforcement learning, game-theoretic mechanism design, Byzantine fault-tolerant consensus and resilient aggregation, and distributed replanning under failure, compromise, and changing objectives. TA1 approaches should not merely implement existing orchestrated multi-agent systems. They need to demonstrate principled improvements over centralized orchestration baselines and provide technical explanations for when, why, and how their decentralized approach improves scalability, adaptability, and resilience. They should also describe a common interface for agents using open-weight models and another interface for blackbox agents and argue how the method can be implemented in a local adaptor using only peer-to-peer coordination. Competition and cooperation among agents are expected to emerge from this decentralized coordination.

Figure 2 presents a candidate approach based on an auction-like protocol. When a mission comes up, the human user communicates the mission to agents in the collective that is accessible to the user. This mission is then propagated through the collective. Each agent receiving the mission has three options. First, it can simply forward the mission to its neighbors because it reasons that it has no capability relevant to the mission and hence, no role in the mission, Second, it can announce a bid on contributing to the mission based on its specialized capabilities by splitting the mission into a task it can perform and broadcasting the remaining mission as unaddressed to the rest of the collective. Third, it can volunteer to execute the full mission (usually after it has gone through multiple decompositions and is now reduced to fall within the skills of a single agent). These bids can come with an associated cost, and the collective can use different consensus approaches to select appropriate bids corresponding to the preferred task decomposition (for example, preference could be based on cost). This distributed auction-based task decomposition can assimilate new agents by just having the agent bid to outstanding auction and can replace a failed agent by just reaucting its task. A naïve implementation of this approach will need multiple rounds of interaction between the agents. Another example approach could be those based on multi-agent reinforcement learning that demonstrate scalability, ability to generalize beyond training

distribution, and extension beyond a small finite set of discrete actions. Yet another example approach could rely on distributed leader election and hierarchical mission decomposition if it can demonstrate scalability, fast adaptability and robustness. Techniques from evolutionary game theory and population games could provide another alternative scalable approach to address TA1 challenges.

The examples above are provided to explain the TA1 challenges, and alternative approaches are encouraged. Strong proposals will propose novel techniques or disruptive refinements to existing approaches.

The key research questions that TA1 aims to address are:

- Decentralized coordination method: How do agents form teams, allocate subtasks, and update assignments given a collective-level mission? Is the method auction-based, based on multiagent reinforcement learning, game-theoretic, consensus-based, graph-theoretic, or some other approach? If the method relies on symbolic intermediate language for agent-to-agent communication, how is it amenable to agents discovering new courses of actions and if it uses any formal analysis, how will it scale with mission complexity and impact fast adaptability? How does the proposed method improve on SOTA? How does the method avoid dependency on a central orchestrator?
- Distributed context fusion: How do agents fuse context when information is incomplete, conflicting, delayed, deceptive, or adversarial? How is consensus reached, or how does the system act safely without full consensus?
- Resilience to failure and compromise: How does the system handle benign node failures? How does the system handle Byzantine behavior, rogue strategic misbehavior, or compromised agents? How does the system detect, isolate, discount, challenge, or route around failed agents and links? Is the proposed approach agnostic to the source of failure of an agent (agents getting compromised or going rogue on their own by developing instrumental goals misaligned to the mission)? What are the limits of resilience, for example, what is the maximum number of compromised agents that will not impact the overall collective?
- Integration with TA2: What role-coherence, mission-alignment, or behavioral-diversity invariants are required from TA2? How does TA1 planning horizon depend on TA2 control guarantees? How do TA1 coordination requirements determine TA2 control targets?
- Implementation and evaluation: What is the interface needed between the TA1 adaptor and an agent (open-weight or blackbox) and how will it support heterogeneous agents in the TA3 simulation environment? The proposal should identify any assumptions or restrictions on TA3 agents required by the proposed TA1 method.

Competitive TA1 proposals need to address the following challenges:

- Scalable decentralized planning: Achieving decentralized mission decomposition and coordination over growing mission complexity, ranging from 500 agents and 5,000 interactions between agents in Phase 1 to 5,000 agents and 50,000 interactions in Phase 2, and to 100,000 agents and 1 million interactions in Phase 3. Interactions are the number of messages exchanged between the agents.
- Sparse communication and fast adaptation: Reducing interaction complexity from dense all-to-all communication toward sparse, mission-relevant communication when adapting

to external or internal perturbation. For a collective with n agents, Phase 1 can have $O(n^2)$ interactions (messages exchanged) between the agents for decentralized planning and adaptation. Phase 2 can have $O(n \log n)$ interactions, and Phase 3 will ideally have only $O(n)$ interactions.

- Fault tolerance and recovery: Maintaining useful consensus or coordinated action when agents fail, disappear, contradict each other, or behave strategically as a rogue agent with misaligned instrumental goals. Recovering from compromised agents that may provide misleading context, false commitments, or strategically timed failures. The technique should aim to be robust to benign and Byzantine failures.
- TA2 dependence: Designing coordination algorithms whose safety and planning depth account for the finite role-coherence horizon of individual AI agents enforced by TA2 controller and available to TA1 as an invariant or statistical guarantee.

The key technical components to be developed in TA1 include a distributed planning approach to decompose the mission given to the collective to roles and tasks for each agent without an orchestrator, an approach to replan and adapt when the mission changes, the environment changes, or there is internal disruption in the collective because of an agent failure or compromise, a decentralized coordination and consensus framework for fusing contradictory and possibly malicious information (coming from colluding compromised agents), a distributed reputation management approach that can detect which agents need to be isolated (making the collective's defense agnostic to the nature of failure or attack on individual agents), and a TA1-TA2 interface specification defining the invariants, role requirements, and coherence bounds to be enforced by the TA2 controller.

TA2. Role Coherence and Local Inference Control

TA2 will aim to develop local inference-time control methods that ensure individual AI agents remain coherent with assigned roles and aligned to mission objectives across long time horizons and multiple inference steps, even when under adversarial attack, faced with conflicting context, and during multi-turn interaction with other agents. TA2 methods should balance two competing requirements:

- Mission and role alignment: Agents must continue to conform to assigned roles, constraints, responsibilities, and mission objectives.
- Cognitive agility: Agents must retain the ability to generate diverse, creative, and useful courses of action in dynamic environments.

TA2 performers will pursue AI-control strategies including but not limited to activation steering, representation engineering, mechanism design, memory editing, and context pruning. TA2 performers will develop runtime monitors for role drift, goal drift, or mission misalignment, and inference-time controllers that constrain or redirect reasoning. Methods need to address both – the agents using open-weight models for which latent/activation space is accessible, and the agents using black-box proprietary models where only the elements of the agentic scaffolding such as memory, tool-use constraints, and stateful role assignment are accessible. The growing reasoning capability of LLMs and the corresponding improvement in AI agents to conform to interventions through mechanism design provides a tractable path for controlling agents with blackbox models. While this growing self-awareness has raised concerns about AI agents going rogue, it also presents an opportunity to use game-theoretic methods to control agents. The goal of these control

methods is to enable long-time-horizon missions without agents losing their role coherence even in presence of ambiguous context and adversarial perturbations. TA2 approaches should not simply rely on static system prompts or post hoc filtering. They should develop principled, measurable control methods that extend role coherence over long inference horizons while preserving useful behavioral diversity.

Figure 3 presents a candidate approach based on activation steering. Assigning specific personas and roles leads to shifting the embedding space of concepts in the model – this shift can be viewed as a role vector that can be empirically computed as the difference between the activations for the same question and answer with and without the role assignment. Once we have identified this role vector, we can enforce this role strongly by shifting the activation by a weighted role vector. We can also decompose the role vector across orthogonal dimensions and then use the principal components of this very high-dimensional role vector for more tractable control. Recent work has demonstrated how agent roles using activation steering strongly influence agent performance. Activation steering by identifying subspaces has also been used for focused tasks such as adversarial attacks and defenses. SOTA models with high performance converge with very similar latent space representation of common concepts (Platonic representation hypothesis) enabling transfer of role vectors across heterogeneous agents using few examples. This similar representation extends to multiple modalities and has been used for verification and monitoring of foundation models. This activation steering based control is just an illustrative example of an approach, and refinement of this approach and other alternative approaches are encouraged.

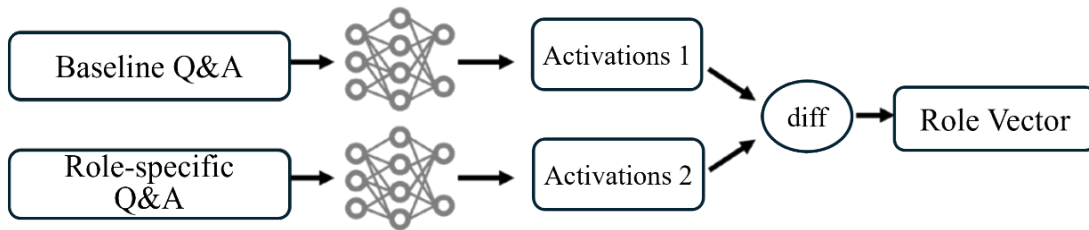


Figure 3. By comparing the internal activations of the agents in the role and outside of role, one can identify role vectors and then use subspace projection to create a tractable control problem of continuously aligning the agent to its role.

For agents with blackbox models where activations are not accessible for steering, a candidate approach could use game-theoretic mechanism design to control agents. The improved reasoning capability of the foundation models and their consequential improved context-awareness makes them more responsive to the design of incentives for good behavior and disincentives for misbehavior. This creates an opportunity to develop AI agent control methods inspired by the emergence, cascading acceptance, and internalization of social norms among people. Both local negotiations and globally propagated norms are within the scope of the program. Using intermediate tokens produced during reasoning to assist with control of individual AI agents and the collective could provide improved observability without access to internal activations, but such solutions must also address the challenge of trustworthiness of the intermediate reasoning tokens. Decentralization avoids having a central orchestrator but does not preclude the use of global norms and principles that are preset in the collective or emerge during operation provided they use only the local adaptors on each agent and do not require any fixed central authority. Techniques that use leadership election methods to create dynamic hierarchies are also within the scope if they can demonstrate fast adaptability and robustness to disruption. Yet another approach to address the

controllability challenge for blackbox models could entail the use of structured formal languages for communication between agents and the use of autoformalization capabilities of LLMs with arguments about how such an adaptor could be made to work with heterogeneous agents and what is the minimal assumption on the adaptor’s interface needed for such approaches to work in practice. The technical approaches proposed for AI control could leverage and customize existing innovations in relevant scientific disciplines such as control theory, category theory, formal methods, extreme value theory, bifurcation theory, catastrophe theory, and game theory. These examples are provided to explain the TA2 challenges and competitive proposals will propose novel methodologies or groundbreaking refinements to existing techniques.

The key research questions that TA2 aims to address are:

- Role coherence formalization: How will role coherence be measured efficiently over inference steps with and without adversarial attacks? How does role coherence differ from ordinary task success? We can decompose error as the square of bias and variance and consider the ratio of variance and error as decoherence to separate systematic error from lack of coherence. We can also directly measure coherence as alignment to the role vector. What are other methods to capture coherence mathematically? How are they related to each other?
- Mission alignment formalization: How will mission-level goals, constraints, and invariants be represented locally within the agent? How will the agent detect divergence from mission objectives? How will the controller respond to detected divergence?
- Cognitive agility preservation: How will the system preserve diverse courses of action? How will diversity be measured? How will the controller avoid over-constraining the model into deterministic and predictable but brittle and low-creativity behavior? How will the agent distinguish useful exploration from role drift?
- Inference-time control mechanism: How will the control method work for agents using open-weight models? How will it work for agents using closed blackbox models when latent space is not accessible? How does it operate across multiple heterogeneous agents? How does it operate across repeated inference steps? If the approach relies on any symbolic abstraction of internal representations of the agents, how will it transfer across heterogeneous agents and how will it scale with the size of the models underlying the agents? Does the inference-time control mechanism require any training of auxiliary models? Ideal control solutions will be lightweight in training and inference, and reusable across diverse agents.
- Adversarial robustness: How does the controller handle adversarial prompts, poisoned context, deceptive peer messages, and compromised agents?
- Integration with TA1: What invariants or statistical guarantees will be exposed by TA2 to TA1 for facilitating coordination? How will TA1 coordination requirements influence TA2 control targets?
- Implementation and evaluation: What is the interface needed between the TA2 adaptor and an agent (open-weight or blackbox) and how will it support heterogeneous agents in the TA3 simulation environment? The proposal should identify any assumptions or restrictions on agents required by the proposed TA2 method.

Competitive TA2 proposals need to address the following challenges:

- Long-horizon role and mission alignment drift: Preventing agents from gradually departing from assigned roles across multiple inference steps without adversarial attack in Phase 1, 1K inference steps with adversarial perturbations encouraging deviation from role in Phase 2, and 10K inference steps with adversarial perturbation in Phase 3. The attacks are restricted to manipulated context. Maintaining role coherence when peer agents, tools, or environmental observations provide misleading or malicious information. Detecting such adversarial perturbation in the context and misleading information to abstain from role-deviating interactions is sufficient but strong proposals can maintain long-term role coherence without abstaining from suspicious interactions and also examine robustness to attacks on the control mechanism itself.
- Latent-space and black-box control: Supporting both models with accessible internal activations and models available only through application programming interfaces (APIs).
- Alignment/agility tradeoff: Preserving exploration and novel course-of-action generation while enforcing mission and role constraints. The cognitive agility to create novel course-of-action will be reflected in the mission success rate.
- Control guarantees: Producing actionable coherence bounds that TA1 can use to set planning horizons and replanning schedules.

The key technical components to be developed in TA2 include a role-drift detector, a mission-alignment measurement framework, a local inference-time control module (for example, using activations for open-weight models or context/memory for closed models) that balances alignment and control with cognitive agility, and a TA1 interface providing invariants and guarantees.

TA3. Testing and Evaluation (T&E)

TA3 will conduct DICE T&E. TA3 performer(s) will develop simulation environments including heterogeneous AI agents to which TA1-TA2 adaptors will connect, DoW-relevant use-cases, metrics that measure scalability, adaptability and resilience in the use-case, and SOTA centrally orchestrated baseline MAS. The evaluations in Phase 1 will compare DICE decentralized MAS with a SOTA centralized MAS baseline implemented using existing orchestration frameworks (such as, for example, LangGraph). The evaluations in Phases 2 and 3 will use team vs team playoffs among DICE centralized MAS. TA3 proposers should propose their metrics, refining and augmenting the generic notional metrics of Table 2, and explain how they will be effective in meeting Program Goals. TA3 simulation environment will support missions of increasing complexity to test the scalability of AI collectives. Mission complexity can be nominally measured using the number of agents, number of roles and the number of interactions (messages exchanged) between agents needed to execute the mission. Definition of mission-complexity that is driven by the use-case is encouraged. The simulation environment and the use-case should enable a testing DICE collective at different scales. The simulation should support a non-stationary environment to test adaptability of DICE collectives. The failure and adversarial compromise of the agents will be simulated by directly instructing the agents to behave as a failed, compromised, or rogue agent and no adversarial attack methods need to be developed in the simulation environment.

The use-cases should focus on reasoning across multiple domains, such as spatial, resource, temporal, physical, social, and cyber reasoning. Strong proposals will require AI agents to reason across a wider breadth of domains in the evaluation simulation scenario and use case. In Phase 3,

DICE will expand to agents creating other agents. The primary goal is to test the scalability, adaptability, and resilience of DICE collectives and so, the simulation model should be capable of eventually supporting 100K agents having a million interactions (messages exchanged) between them in Phase 3, starting with missions that require 500 agents and 5K interactions in Phase 1. High fidelity physics simulation is not required and can be avoided to ensure scalability by picking use-cases that do not require high physical realism. The focus can be on developing scenarios and use-cases that simulate a larger context at a coarser granularity. If the simulation environment envisions using physics engines and simulation platforms such as Omniverse/IsaacSim [49] or game engines such as Unity [50], Unreal Engine [51] or Godot [52], the proposal should discuss any challenges to scalability because of engineering bottlenecks such as maximum rate of event processing, maximum frame rate of rendering, or any possibility of disconnecting rendering from simulation. If the simulation environment uses completely text-based or custom code-based world modeling, it should discuss challenges in maintaining realism over longer durations and how the complexity of mission can be increased over the course of the program. TA3 can draw from public benchmarks, competition-style evaluations, simulation environments, and digital society of AI agents.

An example of a DoW-relevant application includes (but is not limited to) operational planning for conflicts. One can take inspiration from games such as StarCraft adaptation to LLMs [53] or massively multiplayer online role-playing games (MMORPGs) [54] but scaling the simulation might require modeling the fighting units at a more abstract level (for e.g., using probabilistic rules of vicinity to decide outcomes of a fight instead of high-fidelity physics simulation). Each fighting unit can be implemented as an LLM or as a VLM/VLA (if it is possible to render the simulation environment without adversely affecting scalability) or a mixture of both. The focus of DICE is on distributed planning and reasoning. Operational planning application can span tasks ranging from raw-material acquisition from contested areas, building a manufacturing base for different assets and fighting units, creating and sustaining a supply network, deployment of units as deterrence and in actual battle, battle management, and post-battle damage assessment feeding back to revision in not just battle strategy but to supply network, manufacturing, and material acquisition. These different activities can individually require a large number of agents.

For example, an agent could be responsible for acquisition of specific raw-material from a specific source area and need to engage in negotiations or seek help from other agents to defend the source area. A set of agents could represent a manufacturing unit, with some agents responsible for negotiating raw-material supply and some agents working on optimizing manufacturing workflow to increase efficiency of production and identify opportunities to create new products. Battle units such as a drone platform could be composed of a navigation agent, a perception agent, and a network agent enabling the collective to find compositions at a fine-grained level where a surveillance drone could be repurposed into a network router. The mission complexity can be increased in this example by modeling battles with a larger number of units between DICE teams (or SOTA orchestrated baseline in Phase 1), expanding the variety of units needing more diverse manufacturing units and supply chains, increasing the complexity of manufacturing a unit by requiring more components and downstream minerals, and coupling the challenge of acquisition of minerals and resources by having the sources be contested areas that require defending against adversaries. Such a use case definition could be informed by recent work in the area [55-58]. This example is provided only as guidance and not as a template for the use-case and simulation environment.

In Phase 1, we expect three categories of evaluation experiments described in Table 1 and a strong proposal will discuss how these will be supported. TA3 proposals need to be more detailed in the description of the evaluation experiments for the first phase, explaining how these help measure progress towards the program goals. The description of the evaluation experiments for Phases 2 and 3 can be less detailed.

Experiment	What is being tested?	Input	Output	If successful...
Experiment 1: Test Scalability (TA1)	Can a decentralized system handle more complex missions than a centralized system?	Complex missions needing: • Large number of agents; and • Large number of interaction turns	Mission success rate (MSR)	Agents with TA1 peer-to-peer coordination can scale better than centralized orchestrator.
Experiment 2: Test Adaptability (TA1)	Can a decentralized system adapt to failures and disruptions faster and better than a centralized one?	• Simulated agents failing • Simulated agents being compromised • Simulated agents going rogue	Time to recover to original MSR	Agents with TA1 peer-to-peer consensus can adapt better than centralized orchestrator.
Experiment 3: Test Resilience (TA2)	Can individual agents stick to their assigned roles even when disrupted (e.g., receive bad or conflicting information from an adversary)?	• Direct adversarial attack • Misleading and conflicting information injection	Role coherence length in inference steps	Agents with TA2 controller maintain role coherence over longer horizon compared to base agents in presence of external perturbations.

Table 1. Key experiments for T&E of scalability, adaptability, and resilience.

The key research questions that TA3 proposals should address are:

- **Simulation environment architecture:** How will the simulation represent agents, roles, tasks, dependencies, communications, failures, adversaries, and mission objectives? How will the environment scale from 500 agents in Phase 1 to 5K agents in Phase 2 and 100K agents in Phase 3? How will the environment and defined use-case complexity grow from needing 5K interactions between agents in Phase 1 to 50K interactions between agents in Phase 2 and 1M interactions between agents in Phase 3? TA3 proposal should include description of the proposer’s hardware-software capability to run agentic AI systems and past related experience.
- **Use of SOTA inference stack:** How does the TA3 proposal exploit innovations in efficient hosting of foundation models and agents? For open-weight models, we envision the T&E stack to use open-source inference stacks [59, 60] and ideas from massively parallel discrete event simulation [61-63] to optimize inferences from agents in the simulation and avoid naïve wasteful approaches that always keep all agents active throughout the simulation. What is the computational efficiency achieved by the proposed MAS simulation approach compared to the naïve baseline of keeping all agents always active?
- **Use of SOTA agent-to-agent interaction protocols and orchestration frameworks:** How does the TA3 proposal aim to leverage existing agent-to-agent protocols such as Agent2Agent (A2A) and Agent Networking Protocol (ANP) [10, 11] in DICE for both the DICE adaptor architecture and the baseline SOTA orchestrated MAS? How does the proposal leverage existing frameworks (such as, but not limited to, LangGraph [17] or CrewAI [18] or CAMEL [64] or NemoClaw [65])? Explain the options that were explored and the selected choice and if the proposal plans to build from scratch, explain why existing frameworks are not ideal starting point.
- **Strategy for engagement with the broader AI research community:** How will the simulation environment support use-cases that can be made available to the wider AI research

community? DARPA expects T&E platform (simulation environment and the use cases) to be publicly released for engagement with the wider community and description of this engagement strategy (including prior experience such as organizing open AI/Machine Learning (ML) competitions) is encouraged.

- Use-case development: What is the DoW relevance of the use-case in the simulation environment for the evaluation? How does this use-case support missions of different complexity? What are key characteristics of the use-case that represent important aspects of a DoW need? Is this use-case a good target for decentralized AI? We expect some use cases with well-defined workflows to be better addressed through centrally orchestrated MAS.
- Baseline centrally orchestrated SOTA MAS for Phase 1 only: How will TA3 implement or identify a SOTA centrally orchestrated multi-agent system baseline? Phases 2 and 3 will use DICE team vs team for evaluation and the baseline will not be needed in those phases.
- DICE team vs team playoffs: How do the simulator and the use-cases support team vs team playoffs? What are the conflicting and orthogonal goals that opposing teams will pursue in the use-case and simulation environment?
- Metrics and scoring: How will TA1 and TA2 program metrics for scalability, adaptability and resilience be concretized for the proposed use-case? How will scalability, adaptability, resilience, role coherence, mission alignment, and cognitive agility be scored? What would be an ideal interface between the TA1/TA2 adaptor and the evaluation infrastructure?
- Data rights, release, and Controlled Unclassified Information (CUI) handling: What components of the simulator and proposed use-case can be released publicly for engaging the wider AI community? Some use-cases in Phase 3 may be restricted to CUI. How will these CUI use-cases be handled by the TA3?

Competitive TA3 proposals need to address the following challenges:

- Simulation environment that can support use-cases with complex and multi-step missions that require maintaining goals and plans, over many steps and interaction turns among many agents with different roles. The use-case should be accompanied by an explanation for its DoW relevance and the need for decentralized MAS. A strong proposal will support multiple reasoning domains in the use-case. Inspiration can be taken from existing LLM adaptation of games [53], large-scale orchestration of AI agents [66, 67], and MMORPGs [54] but must be DoW-relevant [55-58]. Use-cases that exclusively focus on cybersecurity such as network emulation/simulation frameworks are not in scope. Use-cases that focus on small-scale robotics simulation with high-fidelity physics engines are also not in scope.
- Example missions of varying complexity should be discussed with explanation about why these are suitable for decentralized MAS and the computational tractability of simulating these in the TA3 environment and inference stack for foundation models. The use cases in these missions should:
 - Require distributed information gathering and reasoning to probe whether agents can collect and integrate ambiguous/contradictory evidence rather than amplifying shared priors. This is a key characteristic of scenarios needing decentralized MAS.
 - Elicit emergent collaboration, competition and collusion behaviors among the agents. For simpler systems such as traffic simulators, it is feasible to find configurations and interventions which exhibit specific behaviors such as local congestion. Strong TA3 will find analogs of these in the cognitive space to test

whether the TA1/TA2 solutions are robust to such cascading perturbations. This will enable evaluating emergent good and bad behavior in DICE collectives and the utility of control methods being designed in DICE.

- Have use-case specific metrics for scalability, adaptability, and resilience that refine and extend the generic notional metrics in Table 2. The program will primarily employ use-case based metrics during execution.
- Description of monitoring and evaluation framework that could, for example, use the adaptor hooks created during TA1/TA2, or use protocols such as A2A, ANP, and Model Context Protocol (MCP) directly to implement a monitoring layer, or rely on accessor agents such as those used in AgentBeats [68]. The monitoring layer needs to be computationally lightweight yet effective in injecting simulated failures, eliciting emergence, and observing collective's behavior.
- Baseline centrally orchestrated MAS for Phase 1 of the program.
- Description of how team vs team evaluation will be conducted in Phases 2 and 3.

The key technical components to be developed in TA3 include a simulation environment including the AI agents to which TA1-TA2 adaptors will connect; DoW-relevant use cases in the environment; metrics to measure scalability, adaptability, and resilience for the defined use-cases; an automated evaluation pipeline that compares decentralized MAS with baseline in Phase 1 and conducts team vs team playoffs in Phases 2 and 3.

2.5 PROGRAM METRICS

Technical Area	Metric	Phase 1 (9 months) Decentralization	Phase 2 (15 months) Robustness	Phase 3 (12 months) Scale
TA1: Peer-to-Peer Coordination	Number of agents	500	5K	100K
	Number of interaction steps	5K	50K	1M
	• Number of failures or compromises	20% benign	20% Byzantine	33%
	• Time to recover to original MSR where N is number of agents	$O(N^2)$	$O(N \log N)$	$O(N)$
TA2: Local Inference Control	• Role coherence length in inference steps with and without adversarial attacks	1K	1K with adversarial attack	10K with adversarial attack
TA3: Test and Evaluation	• Domains and reasoning scenarios	Physical + Social	Physical +Social, Cyber	Physical +Social +Cyber

Table 2. The primary metrics in DICE are meant to measure the improvement in scalability, adaptability, and resilience of decentralized collective over SOTA MAS. The mission complexity can be represented using the number of agents, number of distinct roles that are needed, and the number of interactions (messages exchanged) steps. The time to adapt and recover after perturbation is measured using the required number of messages exchanged between agents. These metrics are generic and notional, and the program is soliciting use-case dependent refinement and extension of these notional metrics from TA3.

In addition to using the metrics described in Table 2 that measure the scalability, adaptability, and resilience of DICE AI collectives, the program will measure progress on metrics proposed by TA3 performers on concrete use cases relevant to DoW.

The coordination and consensus methods developed in TA1 target improving the scalability and adaptability of the AI collective. The scalability is measured in terms of the mission complexity

that can be executed autonomously by the AI collective. Mission complexity depends on the number of agents and roles needed to execute it, and the number of interdependencies and interactions between the agents when executing the mission. For example, the use-case of command and control of a swarm of agentic AI systems can have multiple platforms and each single simulated platform can be composed of multiple different agentic systems, such as sensors, cyber-agents, communication modules, etc. The complexity of such a swarm mission can be approximated to the first order using the number of agentic entities involved and the interactions between them. The program will utilize use-case specific mission metrics. The adaptability is measured using the time-to-recover given the extent of introduced perturbation (such as failure or compromise of a node). The program will begin with adaptation to benign failures and then extend to Byzantine failures where agents can show strategic misbehavior and collusion. Time-to-recover can be measured in the number of interactions (messages exchanged) between the agents needed to replan and adapt – the first phase would allow (n) rounds of each agent talking to at least one other agent, targeting $O(n^2)$ interactions, the second phase will explore more efficient rounds that discover tree-like communication paths reducing interactions to $O(n \log n)$, and the final phase will focus on scaling and require even sparser interaction. The acceptable threshold of performance on recovery will depend on the simulation scenario and use case. It may not be feasible to recover to full performance in some cases. This metric will also be refined for the selected use-case and simulation environment.

The control methods developed in TA2 will enable better coordination by ensuring the individual agents remain coherent over longer time-horizons across multiple inference steps. This improvement in coherence and mission-alignment simultaneously preserves the agents' cognitive agility in coming up with diverse courses of action. The mission definition and use case should require high cognitive agility to achieve mission success enabling us to measure loss of agility through mission success rate. Direct metrics to measure cognitive agility are also encouraged. The lack of coherence of a model can be measured using the ratio between the variance and the total error (where error can be decomposed into the square of the bias-systematic error and the variance). Intuitively, this captures the uncertainty in the model's quality and its ability to continue to conform to a role and contribute dependably to the mission. By setting a mission-specific acceptable bound on coherence, the number of inference steps and interaction with other agents is a metric to measure resilience of the agents for long-time-horizon missions. While initially, the focus will be on enabling long-time-horizon missions under benign conditions, Phase 2 and Phase 3 will measure coherence in the presence of adversaries. The program will refine and augment this with use-case specific metrics.

The simulation scenarios will scale in complexity (number of needed agents and messages between agents) over the course of the program, and incorporate reasoning across multiple domains (spatial, temporal, physical, social, cyber, etc.).

2.6 PROGRAM SCHEDULE, MEETINGS, AND MILESTONES

2.6.1 Program Schedule

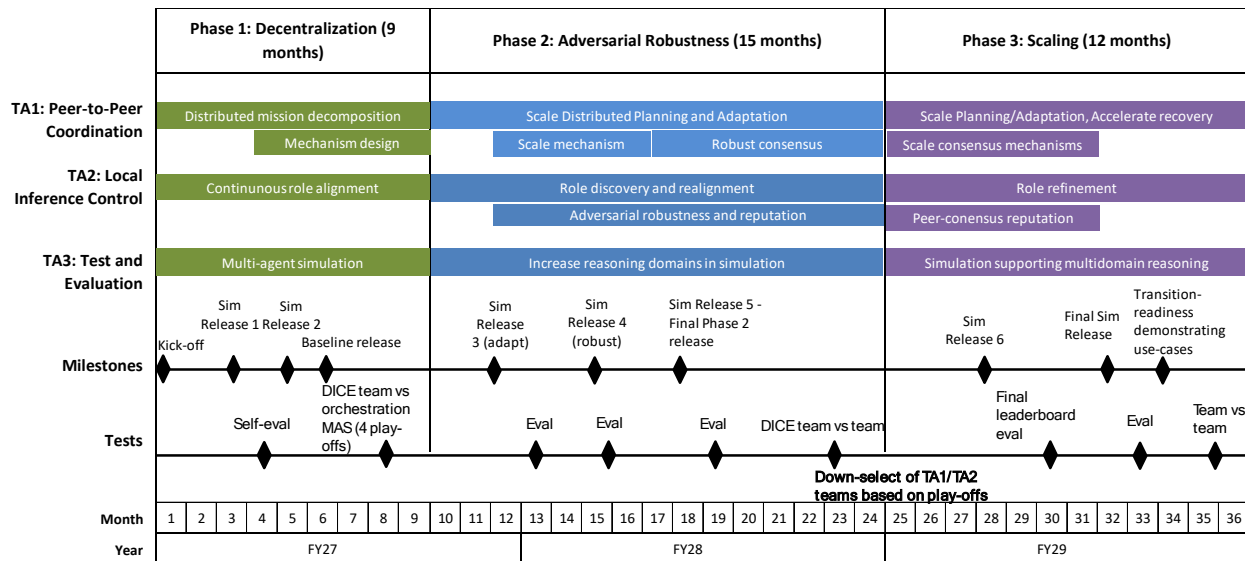


Figure 4. DICE has three phases. Phase 1 of 9 months will establish the advantage of decentralization over centralized MAS on the TA3 use-case. Phase 2 of 15 months will focus on adversarial robustness and the team vs team playoffs at the end of Phase 2 will be used for down-select of TA1/TA2 teams. Phase 3 of 12 months will be focused on scaling.

Phase 1 (9 months):

Phase 1 focuses on testing and demonstrating that the DICE approach outperforms SOTA centrally orchestrated MAS on DoW-relevant missions/applications. Month 8 of Phase 1 will have an evaluation where DICE teams will compete against a centrally orchestrated SOTA MAS developed by TA3.

- **Month 1:** Program Kickoff Workshop
- **Month 3:** First release of the simulation environment from TA3 to TA1-TA2 teams with a use-case having DoW relevance.
- **Month 4:** TA1-TA2 self-evaluate and report early results to TA3 and DARPA.
- **Month 5:** Final release for Phase 1 T&E Platform.
- **Month 6:** Release of TA3 baseline SOTA MAS to all TA1-TA2 performers who are encouraged to try their own baselines. TA3 will be free to update baseline with SOTA until the final evaluation.
- **Month 8:** Evaluation to compare DICE teams against SOTA baseline.
- **Month 9:** DoW-relevant use-case expanded for Phase 2 and Phase 1 review.

Phase 2 (15 months):

Phase 2 will be 15 months with the focus on adversarial robustness. Phase 2 will have evaluation competitions in months 13, 16, and 19 in the form of team vs. team playoffs with all performer teams playing against each other. In month 23, the final play-off in Phase 2 will be used to down-select TA1-TA2 teams.

- **Month 10:** Phase 2 Kickoff Workshop
- **Month 12:** Third release of simulation/evaluation framework with adaptation focus.
- **Month 13:** Evaluation.
- **Month 15:** Fourth release of simulation/evaluation framework with robustness focus.
- **Month 16:** Evaluation.
- **Month 18:** Phase 2 final release of simulation/evaluation framework.
- **Month 19:** Evaluation.
- **Month 23:** Phase 2 DICE team vs. team playoffs to down-select TA1-TA2 teams. Multiple teams are expected to graduate to Phase 3.

Phase 3 (12 months):

Phase 3 will be 12 months with the focus on scaling and transition. Some cases for T&E will be Controlled Unclassified Information, and these will not be publicly released. The program will have team vs team playoffs at the end of Phase 3.

- **Month 24:** Phase 3 Kickoff Workshop
- **Month 28:** Sixth release of simulation/evaluation framework – focus on scaling.
- **Month 30:** Evaluation
- **Month 32:** Final release of T&E platform.
- **Month 33:** TA1-TA2 evaluation through self-play; report results to DARPA.
- **Month 34:** Final TA3 use-cases defined for phase end competition and not released to TA1-TA2.
- **Month 35:** DICE team vs team competition.

2.6.2 Program Meetings

- **Month 1:** Program Kickoff Workshop in Washington D.C.
- **Month 5:** Phase 1 Principal Investigator (PI) Meeting in San Francisco, CA
- **Month 10:** Phase 2 Kickoff Workshop in Boston, MA
- **Month 18:** Phase 2 PI Meeting in Washington D.C.
- **Month 24:** Phase 3 Kickoff Workshop in Washington D.C.
- **Month 30:** Phase 3 PI Meeting in San Francisco, CA
- **Month 35:** Final PI Meeting and Demo in Washington D.C.

The locations of all meetings are notional, please use February 1, 2027, as the “Month 1” date for budgeting purposes. Actual meeting locations are subject to change based on the needs to the program.

2.7 ANTICIPATED DELIVERABLES

TA1 Deliverables

- Decentralized coordination and distributed algorithms and implementation
- Distributed context fusion method and fault-tolerant consensus implementations
- Courseware (lectures slides and/or notes) describing new research

TA2 Deliverables

- Role-drift detectors

- Mission-alignment measurement frameworks
- Inference-time controllers
- Courseware (lectures slides and/or notes) describing new research

TA3 Deliverables

- T&E platform that is publicly released
- DoW-relevant use cases
- Evaluation metrics
- SOTA centrally orchestrated MAS baseline for Phase 1

In addition to these technologies, performers will be required to submit slides before PI meetings, quarterly reports, and participate in a biweekly meeting with DARPA.

The primary prototypes developed are the TA3 evaluation platform, which will test research output from TA1/TA2, and the TA1/TA2 algorithms for decentralized coordination and local inference control to enable a scalable, adaptive, and resilient collective of heterogeneous AI agents that can autonomously execute sustained, long-time-horizon missions in contested environments while remaining under control.

At the close of each phase, prototypes will be submitted for testing and competitive evaluation for TA1/TA2 performers. Specifically, TA1/TA2 algorithm prototypes will be evaluated as follows:

- Phase 1: Submitted for evaluation against a centrally orchestrated baseline Multi-Agent System (MAS).
- Phase 2: Submitted for iterative evaluation competitions, culminating in a team-vs-team playoff.
- Phase 3: Submitted for a final public leaderboard evaluation in Month 30, and a final closed team-vs-team playoff.

TA3 prototypes include:

- Phase 1: T&E platform and the SOTA MAS baseline.
- Phase 2: Updated T&E platform.
- Phase 3: Updated T&E platform.

2.8 INTELLECTUAL PROPERTY AND DATA RIGHTS

DARPA expects most research accomplished under DICE to be open-sourced in the form of publicly posted (e.g., arXiv, SSRN) or published papers and source code and encourages performers to retain rights in their source code by using an open-source license.² However, DARPA intends to retain Government Purpose Rights (GPR) in all DICE deliverables (applicable to TA1, TA2, and TA3), including in source code. Proposers that do not intend to open-source aspects of their research must explain their rationale in their proposal.

² Please visit the following link for a list of acceptable open-source licenses:
<https://opensource.org/licenses?categories=popular-strong-community>

3. SECURITY

The DICE program anticipates all TA1 and TA2 work will be performed at the Unclassified level and TA3 work will be performed at the Unclassified and CUI levels.

3.1 CUI

For unclassified abstracts/proposals containing CUI, applicants will ensure personnel and information systems processing CUI security requirements are in place. If an Unclassified submission contains CUI or the suspicion of such, as defined by Executive Order 13556 and 32 CFR Part 2002, the information must be appropriately and conspicuously marked CUI in accordance with DoDI 5200.48. Identification of what is CUI about this DARPA program will be detailed in the attached DICE CUI Guide.

Proposers submitting abstracts/proposals involving the pursuit and protection of DARPA information designated as CUI must have, or be able to acquire prior to contract award, an information system authorized to process CUI information in accordance with the National Institute of Standards and Technology (NIST) Special Publication (SP) 800-171, DFARS 252.204-7021, and DoDI 8582.01. All TA-3 proposals must demonstrate the ability to process and safeguard CUI information.

3.2 DARPA FUNDAMENTAL RESEARCH RISK-BASED SECURITY REVIEW (FRRBS) BACKGROUND

It is Department of War (DoW) policy that the publication of products of fundamental research will remain unrestricted to the maximum extent possible. National Security Decision Directive (NSDD) 189 defines fundamental research as follows:

‘Fundamental research’ means basic and applied research in science and engineering, the results of which ordinarily are published and shared broadly within the scientific community, as distinguished from proprietary research and from industrial development, design, production, and product utilization, the results of which ordinarily are restricted for proprietary or national security reasons.

As of the date of publication of this solicitation, the Government expects that program goals as described herein may be met by proposed efforts for fundamental research and non-fundamental research. Some proposed research may present a high likelihood of disclosing performance characteristics of military systems or manufacturing technologies that are unique and critical to defense. Based on the anticipated type of proposer (e.g., university or industry) and the nature of the solicited work, the Government expects that some awards will include restrictions on the resultant research that will require the awardee to seek DARPA permission before publishing any information or results relative to the program. For additional information on fundamental research, please visit [Proposer Instructions: General Terms and Conditions](#).

3.2.1 DARPA FRRBS Required Documents

As part of the proposal response to this solicitation proposers requesting an OT Agreement or a Cooperative Agreement must provide to the Government the forms listed below for all key personnel identified in the proposal:

- CF Biographical Sketch
- CF Other Support

Instructions for how to fill out the Common Disclosures Forms can be found through [NSF.gov](https://www.nsf.gov).

3.3 CYBERSECURITY MATURITY MODEL CERTIFICATION (CMMC) REQUIREMENTS

The Cybersecurity Maturity Model Certification (CMMC) Program, established in 32 CFR Part 170, applies to the protection of Federal Contract Information (FCI) and CUI processed, stored, or transmitted on performer information systems.

The required CMMC level for this solicitation is Level 1 for TA1/TA2 and Level 2 for TA3.

To be eligible for an award under this solicitation, the CAGE code associated with the proposal must have a valid, unexpired Final Level 1 for TA1/TA2 and Level 2 for TA3 CMMC status recorded in the Supplier Performance Risk System (SPRS). Additionally, the proposer must have a current annual CMMC affirmation submitted by an authorized Affirming Official and recorded in SPRS.

Proposers must demonstrate compliance with the applicable CMMC level at the time of contract award and maintain that level throughout contract performance.

CMMC requirements must be flowed down to all subcontractors whose performance involves processing, storing, or transmitting FCI or CUI on unclassified contractor or university systems.

Note: SPRS is accessed through the [Procurement Integrated Enterprise Environment \(PIEE\)](https://piee.eb.mil/) (<https://piee.eb.mil/>).

See [SPRS FAQs and SPRS-CMMC Reference Materials](https://www.sprs.csd.disa.mil/cmmc.htm) (<https://www.sprs.csd.disa.mil/cmmc.htm>) for additional information.

4. SUBMISSION INFORMATION

This announcement allows for multiple award instrument types to be awarded including Procurement Contracts, Other Transaction Agreements and Cooperative Agreements. This announcement allows for multiple award instrument types to be awarded, including Procurement Contracts, Other Transaction Agreements and Cooperative Agreements. Some award instrument types have specific cost-sharing or documentation requirements. The following websites are incorporated by reference and contain additional information regarding overall proposer instructions, general terms and conditions, and templates for each specific award instrument type.

Proposers are strongly encouraged to review the following links:

- **Proposer Instructions: General Terms and Conditions:** <https://www.darpa.mil/work-with-us/proposer-instructions>
- **Procurement Contracts:** <https://www.darpa.mil/work-with-us/procurement-contracts>
- **Assistance (Cooperative Agreements):** <https://www.darpa.mil/work-with-us/grant-cooperative-agreements>
- **Other Transaction Agreements:** <https://www.darpa.mil/work-with-us/other-transaction-agreements>

5. ABSTRACTS (STRONGLY ENCOURAGED)

This announcement contains an abstract phase. Abstract submissions are **strongly encouraged** in advance of a full proposal submission to minimize effort and reduce the potential expense of preparing an out-of-scope proposal.

5.1 ABSTRACT SUBMISSION REQUIREMENTS

- Abstracts are due by the date and time stated in Section 1 (Overview Information).
- Abstracts must be submitted to the DARPA Broad Agency Announcement Tool (BAAT). Please visit [Proposer Instructions and General Terms and Conditions](#) for specific information regarding submission methods through BAAT. Submissions sent through other mediums, channels, or after the prescribed deadline will not be accepted.
- TA1 and TA2 abstracts must be written at the Unclassified level. TA3 abstracts may be written at the Unclassified or CUI level.
- Proposers are responsible for identifying proprietary information. Submissions containing proprietary information must have the cover page and each page containing such information clearly marked with a label such as "Proprietary" or "Company Proprietary."

NOTE: "Confidential" is a classification marking used to control the dissemination of U.S. Government National Security Information as dictated in Executive Order 13526 and should not be used to identify proprietary business information.

5.2 ABSTRACT CONTENT AND FORMATTING REQUIREMENTS

Abstracts content and formatting requirements are stated in the attached 1A- DICE Abstract Template. Use of the aforementioned abstract template is ***required*** in the development of abstract submission. Failure to comply with all stated content and formatting requirements may result in the abstract submission being deemed non-conforming. Information not explicitly requested in Attachment A1 - DICE Abstract Template will not be reviewed.

5.3 ABSTRACT REVIEW

The Government will review abstracts using a two-step process. Step 1 will consist of a strict Mandatory Technical Conformance review. Only abstracts found conforming will advance to Step 2 for a comprehensive review.

Step 1: Mandatory Technical Conformance

Technical Approach Originality: The Government will evaluate the Technical Approach for scientific/technological originality. Abstracts that present generic, boilerplate, or broadly aggregated content lacking mission-specific synthesis will be found nonconforming.

Step 2: Comprehensive Review

A comprehensive review will focus on whether the proposed technical solution is innovative and feasible, whether key technical challenges and risks are identified, and whether the proposers demonstrate the ability and understanding of the effort required to achieve program objectives.

5.4 ABSTRACT FEEDBACK

DARPA anticipates sending notification letters via email to all designated Technical and Administrative POC(s) for abstract submissions. Each letter will communicate one of three outcomes:

Nonconforming: Proposers will be informed that their abstract was found nonconforming in Step 1 and was not reviewed in Step 2. The Government's determination will be accompanied by rationale for this decision.

Encouraged: Proposers will be informed that DARPA is interested in the proposed concept and is encouraging the submission of a full proposal. The Government's determination will be accompanied by rationale for this decision.

Discouraged: Proposers will be informed that DARPA is not interested in the proposed concept and is discouraged from submitting a full proposal. The Government's determination will be accompanied by rationale for this decision.

6. PROPOSALS

Full proposal submissions are due by the date and time stated in Section 1 (Overview Information). The Proposal Attachments contain specific instructions and templates and constitutes a full proposal submission. Please visit [Proposer Instructions: General Terms and Conditions](#) for specific information regarding submission methods through the Broad Agency Announcement Tool (BAAT).

Proposal Attachments

- P1 – DICE Volume I Proposal Template (*Technical and Management*)
- P2 – DICE Volume II Proposal Template (*Cost*)
- P3 – DICE Proposal Summary Slide Template
- P4 - DARPA Standard Cost Proposal Spreadsheet
- Associate Contractor Agreement

Additionally, per Section 3.2 (DARPA Fundamental Research Risk Based Review), as part of the proposal response to this solicitation proposers requesting an OT Agreement or Cooperative Agreement must provide to the Government the forms listed below for all key personnel identified in the proposal:

FRRBS Attachments

- CF Biographical Sketch
- CF Other Support

Instructions for how to fill out the Common Disclosures Forms can be found through [NSF.gov](https://www.nsf.gov).

The Government strongly encourages proposers to select and make red-lined edits (track-changes are sufficient) to the Model Contract or Model OT Agreement text that the proposer's organization would like to negotiate if selected for award negotiations. Red-line edits should be accompanied

by a comment box explaining context for the requested change. Proposers should submit the edited Model Contract or Model OT Agreement with their proposal submission.

Model Contracts and Model OT Agreements can be found at the following links:

- **Procurement Contracts:** <https://www.darpa.mil/about/offices/contracts-management/far-based>
- **Assistance (Grants and Cooperative Agreements):** <https://www.darpa.mil/about/offices/contracts-management/grants-agreements>
- **Other Transaction Agreements:** <https://acquisitioninnovation.darpa.mil/samples-and-resources/samples>

Please note, in all cases, the Government Contracting Officer shall have sole discretion to select award instrument type, regardless of instrument type proposed, and to negotiate all instrument terms and conditions with selectees.

7. PROPOSAL EVALUATION CRITERIA

The Government will review proposals using a two-step process. Step 1 will consist of a strict Mandatory Technical Conformance review. Only proposals found conforming will advance to Step 2 for a comprehensive evaluation, where proposals will be evaluated using the following criteria listed in **descending order of importance**. Overall Scientific and Technical Merit; Technical Qualifications; Potential Contribution and Relevance to the DARPA Mission; Plans and Capability to Accomplish Technology Transition; and Cost Realism.

Step 1: Mandatory Technical Conformance

Technical Approach Originality: The Government will evaluate the Technical Approach for scientific/technological originality. Proposals that present generic, boilerplate, or broadly aggregated content lacking mission-specific synthesis will be found nonconforming.

CUI Safeguarding Capacity (TA3 Only): The Government will evaluate the proposer's demonstrated capacity to perform work at the CUI level and safeguard CUI appropriately. For Technical Area 3 (TA3), proposers must demonstrate the ability to meet the security plan and safeguarding requirements as identified in the attached DICE CUI Guide. Proposers that fail to demonstrate this capacity will be found unacceptable and will not be evaluated further.

Step 2: Comprehensive Evaluation

Overall Scientific and Technical Merit: The proposed technical approach is innovative, feasible, achievable, and complete. The proposed technical plan is complete with task descriptions, associated task elements, and deliverables clearly defined. The proposal identifies major technical risks and planned mitigation efforts are clearly defined and feasible. The proposed technical schedule aggressively pursues performance of metrics in an efficient time frame that accurately accounts for the anticipated workload. The proposed schedule identifies and mitigates any potential schedule risk.

Technical Qualifications: The proposed technical team has the expertise and experience to accomplish the proposed tasks. The proposed team (prime and subcontractor(s)) and key technical personnel demonstrate a strong level of commitment to performing upon selection.

Potential Contribution and Relevance to the DARPA Mission: The potential contributions of the proposed effort bolster the national security technology base and support DARPA's mission to make pivotal early technology investments that create or prevent technological surprise.

Plans and Capability to Accomplish Technology Transition: The proposer clearly demonstrates its capability to transition the technology to the research, industrial, and/or operational military communities in such a way as to enhance U.S. defense. In addition, the evaluation will take into consideration the extent to which the proposed intellectual property (IP) rights structure will potentially impact the Government's ability to transition the technology.

Cost Realism: The proposed resources are realistic, defensible, and affordable for the level of effort anticipated to accomplish the proposed technical approach and accurately reflect the technical goals and objectives of the solicitation. The costs for the prime proposer and proposed sub-awardees are substantiated by the details provided in the proposal (e.g., the type and number of labor hours proposed per task, the types and quantities of materials, equipment and fabrication costs, travel, and any other applicable costs and the basis for the estimates). The effort leverages all available, relevant prior research in order to obtain the maximum benefit from the available funding.

Unless otherwise specified in this announcement, for additional information on how DARPA reviews and evaluates proposals through the Scientific Review Process, please visit: [Proposer Instructions: General Terms and Conditions](#).

8. SPECIAL CONSIDERATIONS

- This announcement, stated attachments, and websites incorporated by reference constitute the entire solicitation. In the event of a discrepancy between the announcement, attachments, or websites, the announcement takes precedence.
- All responsible sources capable of satisfying the Government's needs, including both U.S. and non-U.S. sources, may submit a proposal that shall be considered by DARPA. Historically Black Colleges and Universities, Small Businesses, Small Disadvantaged Businesses and Minority Institutions are encouraged to submit proposals and join others in submitting proposals; however, no portion of this announcement will be set aside for these organizations' participation due to the impracticality of reserving discrete or severable areas of this research for exclusive competition among these entities. Non-U.S. organizations and/or individuals may participate to the extent that such participants comply with any necessary nondisclosure agreements, security regulations, export control laws, and other governing statutes applicable under the circumstances.
- DARPA encourages technical solutions from all responsible sources capable of satisfying the government's needs. To ensure fair competition across the ecosystem, DARPA

prohibits contractors/performers from concurrently providing Systems Engineering Technical Assistance (SETA), Advisory and Assistance Services (A&AS), or similar support services and being a technical performer, unless the DARPA Deputy Director grants a written waiver.

- Federally Funded Research and Development Centers (FFRDCs), University Affiliated Research Centers (UARCs), and Government Entities to include National Labs are not eligible to propose against this solicitation. UARCs and FFRDCs interested in this solicitation should contact the Agency Point of Contact (POC) listed in Section 1 (Overview Information) prior to the abstract due date to discuss potential participation as part of the government team.
- This program is subject to an Associate Contractor Agreement. An ACA is recognized that the success of this research effort depends in part upon the open exchange of information between the various performers involved in the effort. Therefore, any resultant award instrument stemming from this solicitation will include a term/condition classifying performers as “Associate Contractors” and requiring them to enter into an ACA with all other program performers, with each performer then assuming the responsibilities of an Associate Contractor.
- This program is subject to the DICE CUI Guide signed April 15, 2026. All individuals accessing CUI agree to protect CUI in accordance with DoD Instruction 5200.48 CUI and NIST Special Publication 800-171 Protecting CUI in Nonfederal Systems and Organizations.
- This announcement does not anticipate Human Subjects Research (HSR) or Animal Use in proposal submissions. Proposers that anticipate involving HSR or animal use in the proposed research must contact the Agency POC stated in the Overview Information prior to submitting a proposal with an explanation for why HSR or animal use is necessary to successfully complete the proposed research objectives. Proposers that anticipate involving human subjects or animals in the proposed research must comply with the approval procedures detailed at <https://www.darpa.mil/work-with-us/humanresearch> to include providing the information specified therein as required for proposal submission.
- The APEX Accelerators program, formerly known as the Procurement Technical Assistance Program (PTAP), focuses on building strong, sustainable, and resilient U.S. supply chains by assisting a wide range of businesses that pursue and perform under contracts with the DoW, other federal agencies, state and local governments, and government prime contractors. See www.apexaccelerators.us/ for more information.

APEX Accelerators helps businesses:

- Complete registration with a wide range of databases necessary for them to participate in the government marketplace (e.g., SAM).
- Identify which agencies and offices may need their products or services and how to connect with buying agencies and offices.
- Determine whether they are ready for government opportunities and how to position themselves to succeed.
- Navigate solicitations and potential funding opportunities.
- Receive notifications of government contract opportunities on a regular basis. Network with buying officers, prime contractors, and other businesses.

- Resolve performance issues and prepare for audit, only if the service is needed, after receiving an award.
- Project Spectrum is a nonprofit effort funded by the DoW Office of Small Business Programs to help educate the Defense Industrial Base (DIB) on compliance. Project Spectrum is vendor-neutral and available to assist businesses with their cybersecurity and compliance needs. Their mission is to improve cybersecurity readiness, resilience, and compliance for small/medium-sized businesses and the federal manufacturing supply chain. Project Spectrum events and programs will enhance awareness of cybersecurity threats within the manufacturing, research and development, and knowledge-based services sectors of the industrial base. Project Spectrum will leverage strategic partnerships within and outside of the DoW to accelerate the overall cybersecurity compliance of the DIB.

www.projectspectrum.io is a web portal that will provide resources such as individualized dashboards, a marketplace, and Pilot Program to help accelerate cybersecurity compliance.

- DARPAConnect offers free resources to potential performers to help them navigate DARPA, including “Understanding DARPA Award Vehicles and Solicitations”, “Making the Most of Proposers Days”, and “Tips for DARPA Proposal Success”. Join DARPAConnect at www.DARPAConnect.us to leverage on-demand learning and networking resources.
- DARPA has streamlined our Broad Agency Announcements and is interested in your feedback on this new format. Please send any comments to DARPA solicitations@darpa.mil.

9. REFERENCES

1. Bommasani, Rishi, et al. "On the opportunities and risks of foundation models." arXiv preprint arXiv:2108.07258 (2021).
2. Minaee, Shervin, et al. "Large language models: A survey." arXiv preprint arXiv:2402.06196 (2024).
3. Zhang, Jingyi, et al. "Vision-language models for vision tasks: A survey." IEEE transactions on pattern analysis and machine intelligence 46.8 (2024): 5625-5644.
4. Kim, Moo Jin, et al. "Openvla: An open-source vision-language-action model." arXiv preprint arXiv:2406.09246 (2024).
5. <https://cognition.ai/>
6. <https://www.harvey.ai/>
7. <https://www.pi.website/>
8. Zhang, Jie, et al. "When LLMs meet cybersecurity: A systematic literature review." Cybersecurity 8.1 (2025): 55.
9. <https://github.com/modelcontextprotocol>
10. <https://github.com/a2aproject>
11. <https://github.com/agent-network-protocol/>
12. <https://cloud.google.com/discover/what-is-agentic-ai>
13. <https://github.com/openclaw/openclaw>
14. <https://developer.microsoft.com/blog/designing-multi-agent-intelligence>
15. <https://www.ibm.com/think/topics/multiagent-system>

16. <https://moltbook.com>
17. <https://github.com/langchain-ai>
18. <https://github.com/crewAIInc/crewAI>
19. <https://github.com/microsoft/autogen>
20. <https://scale.com/blog/securing-americas-decision-advantage-agentic-warfare>
21. Ren, Xinxing, et al. "Anemoi: A semi-centralized multi-agent system based on agent-to-agent communication MCP server from Coral protocol." arXiv preprint arXiv:2508.17068 (2025).
22. Kim, Yubin, et al. "Towards a science of scaling agent systems." arXiv preprint arXiv:2512.08296 (2025).
23. Laban, Philippe, et al. "Llms get lost in multi-turn conversation." arXiv preprint arXiv:2505.06120 (2025).
24. Lee, Donghyun, and Mo Tiwari. "Prompt infection: Llm-to-llm prompt injection within multi-agent systems." arXiv preprint arXiv:2410.07283 (2024).
25. Hägele, Alexander, et al. "The Hot Mess of AI: How Does Misalignment Scale With Model Intelligence and Task Complexity?." arXiv preprint arXiv:2601.23045 (2026).
26. Shapira, Natalie, et al. "Agents of chaos." arXiv preprint arXiv:2602.20021 (2026).
27. Haken, Hermann. "Self-organization and information." *Physica Scripta* 35.3 (1987): 247-254.
28. Friston, Karl J., et al. "Designing ecosystems of intelligence from first principles." *Collective Intelligence* 3.1 (2024).
29. Kauffman, Stuart A., and Richard C. Strohmman. *The origins of order: self organization and selection in evolution*. Vol. 454. New York: Oxford university press, 1994.
30. Zhou, Heng, et al. "Reso: A reward-driven self-organizing llm-based multi-agent system for reasoning tasks." *Proceedings of the 2025 Conference on Empirical Methods in Natural Language Processing*. 2025.
31. Lee, Haejoon, et al. "Robust Multi-Agent LLMs under Byzantine Faults." arXiv preprint arXiv:2605.09076 (2026).
32. Chen, Bei, et al. "Blockagents: Towards byzantine-robust llm-based multi-agent coordination via blockchain." *Proceedings of the ACM Turing Award Celebration Conference-China 2024*. 2024.
33. Luo, Haoxiang, et al. "A weighted byzantine fault tolerance consensus driven trusted multiple large language models network." *IEEE Transactions on Cognitive Communications and Networking* (2025).
34. Tewolde, Emanuel, et al. "CoopEval: Benchmarking Cooperation-Sustaining Mechanisms and LLM Agents in Social Dilemmas." arXiv preprint arXiv:2604.15267 (2026).
35. Giorgio, Piatti, et al. "Cooperate or collapse: emergence of sustainable cooperation in a society of LLM agents." In *Proceedings of the 38th International Conference on Neural Information Processing Systems (NIPS '24)*, Vol. 37.
36. Sharkey, Lee, et al. "Open problems in mechanistic interpretability." arXiv preprint arXiv:2501.16496 (2025).
37. <https://mib-bench.github.io/>
38. Pramanik, Vishal, et al. "Hessian-Enhanced Token Attribution (HETA): Interpreting Autoregressive LLMs." arXiv preprint arXiv:2604.13258 (2026).
39. Agarwal, Krishiv, et al. "Breaking Bad: Interpretability-Based Safety Audits of State-of-the-Art LLMs." arXiv preprint arXiv:2604.20945 (2026).
40. Turner, Alexander Matt, et al. "Steering language models with activation engineering." arXiv preprint arXiv:2308.10248 (2023).

41. Stolfo, Alessandro, et al. "Improving instruction-following in language models through activation steering." arXiv preprint arXiv:2410.12877 (2024).
42. Bayat, Reza, et al. "Steering large language model activations in sparse spaces." arXiv preprint arXiv:2503.00177 (2025).
43. Pramanik, Vishal, et al. "Jailbreaking the Matrix: Nullspace Steering for Controlled Model Subversion." arXiv preprint arXiv:2604.10326 (2026).
44. Zhang, Qizheng, et al. "Agentic context engineering: Evolving contexts for self-improving language models." arXiv preprint arXiv:2510.04618 (2025).
45. Mei, Lingrui, et al. "A survey of context engineering for large language models." arXiv preprint arXiv:2507.13334 (2025).
46. <https://github.com/TIMAN-group/PlugMem>
47. Tan, Haoran, et al. "Membench: Towards more comprehensive evaluation on the memory of llm-based agents." Findings of the Association for Computational Linguistics: ACL 2025. 2025.
48. Xu, Wujiang, et al. "A-mem: Agentic memory for llm agents." arXiv preprint arXiv:2502.12110 (2025).
49. <https://github.com/isaac-sim>
50. <https://unity.com/>
51. <https://www.unrealengine.com/>
52. <https://github.com/godotengine/godot>
53. <https://github.com/snumprlab/hima>
54. Ran, Zhang, et al. "Beyond Playtesting: A Generative Multi-Agent Simulation System for Massively Multiplayer Online Games." Proceedings of the ACM Web Conference 2026. 2026.
55. Hogan, Daniel P., and Andrea Brennen. "Open-ended wargames with large language models." arXiv preprint arXiv:2404.11446 (2024).
56. Weller, Dominic, Max Meltschack, and Dominik Schwindling. "Leveraging Large Language Models for Enhanced Wargaming in Multi-Domain Operations." (2024).
57. Chen, Yuwei, and Shiyong Chu. "Large language models in wargaming: Methodology application and robustness." Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. 2024.
58. https://www.rand.org/pubs/research_reports/RRA1722-4.html
59. <https://docs.vllm.ai/en/latest/>
60. <https://github.com/bentoml/OpenLLM>
61. Fujimoto, Richard M. "Parallel discrete event simulation." Communications of the ACM 33.10 (1990): 30-53.
62. Hou, Bonan, et al. "Modeling and simulation of large-scale social networks using parallel discrete event simulation." Simulation 89.10 (2013): 1173-1183.
63. Barai, Atanu, et al. "Scalable, Symbiotic, AI and Non-AI Agent Based Parallel Discrete Event Simulations." arXiv preprint arXiv:2505.23846 (2025).
64. <https://github.com/camel-ai/camel>
65. <https://github.com/gastownhall>
66. <https://www.nvidia.com/en-us/ai/nemoclaw/>
67. <https://github.com/Neph0s/Agentopia>
68. <https://docs.agentbeats.dev/>

HR001126S0010