

Grounded Reasoning and Artificial Intelligence for National Security (GRAINS) Workshop

Breakout Questions

Objective: The GRAINS Workshop brings together experts to advance the use of formal methods in rigorously ensuring the security AI-enabled national security systems and critical infrastructure. The goal is to foster collaboration, identify challenges and success stories, and develop actionable plans with clear ownership to drive innovation, inform policy, and expand programmatic opportunities.

Background: Modern AI (e.g., agentic LLMs) is nondeterministic. Formal methods are deterministic and can potentially provide guaranteed boundaries on AI systems, especially in national security and critical infrastructure applications. In the workshop we want to consider the fast evolving capabilities in AI for formal methods and formal methods for AI – both affect challenges in assured security and resilience of AI in national security and critical infrastructure applications.

Question & Context	Who do you Trust? Thinking broadly, about the whole ecosystem of national security technology challenges, who/what do you trust for building rigorous security guarantees?	Who is your Adversary? Thinking broadly, about the whole ecosystem of national security technology challenges, who/what thwarts/obstructs building secure capabilities?	What are viable AI-OT Frameworks? For national security (NS) and critical cyber-physical infrastructure (CI), how viable are current and emerging frameworks?	How do Architectures handle AI failures? For national security (NS) and critical cyber-physical infrastructure (CI) applications, what architectures are robust to AI failures of varied types. How is the considerable capabilities of AI harnessed while mitigating the dangers?
Information sought	- Relative/ranked trust of entities, - How has who your trust changed since 2022 or is changing now - What priority actions can we consider to improve trust for specific entities? - Criteria of success	- Relative/ranked impact/concern of adversary entities, prioritize threats - How has the threat landscape changed? - How has who your sense of adversary changed since 2022 or is changing now - What priority actions can we consider to better thwart specific entities? - Near-term actions that can mitigate threats - Criteria of success	- Pre-2022 frameworks that are succeeding or failing for robust incorporation of modern AI into OT/ICS environments. - How has technical debt of existing frameworks changed since 2022 - New/evolved/promising AI-OT/ICS frameworks - What are priority actions needed to advance/mature/secure AI-OT/ICS frameworks for NS/CI applications? - Criteria of success	- Pre-2022 architectures that are succeeding or failing for robust architectures incorporating modern AI into OT/ICS environments. - New/evolved/promising AI-failure resilient frameworks - What are priority actions needed to advance/mature AI-resilient architectures for NS/CI applications? - Criteria of success
Criteria	- SMART actions - Supporting evidence that exists or needed or anticipated - Emerging or on/over the horizon entities/issues for trust	- SMART actions - Supporting evidence that exists or needed or anticipated - Emerging or on/over the horizon entities/issues wrt adversaries	- SMART actions - Supporting evidence that exists or needed or anticipated - Emerging or on/over the horizon AI-OT/ICS frameworks	- SMART actions - Supporting evidence that exists or needed or anticipated - Emerging or on/over the horizon AI-failure resilient architectures for NS/CI applications
Guidance for speakers	Evocative, forward looking, property focused	Evocative, forward looking, property focused	Evocative, forward looking, property focused	Evocative, forward looking, property focused