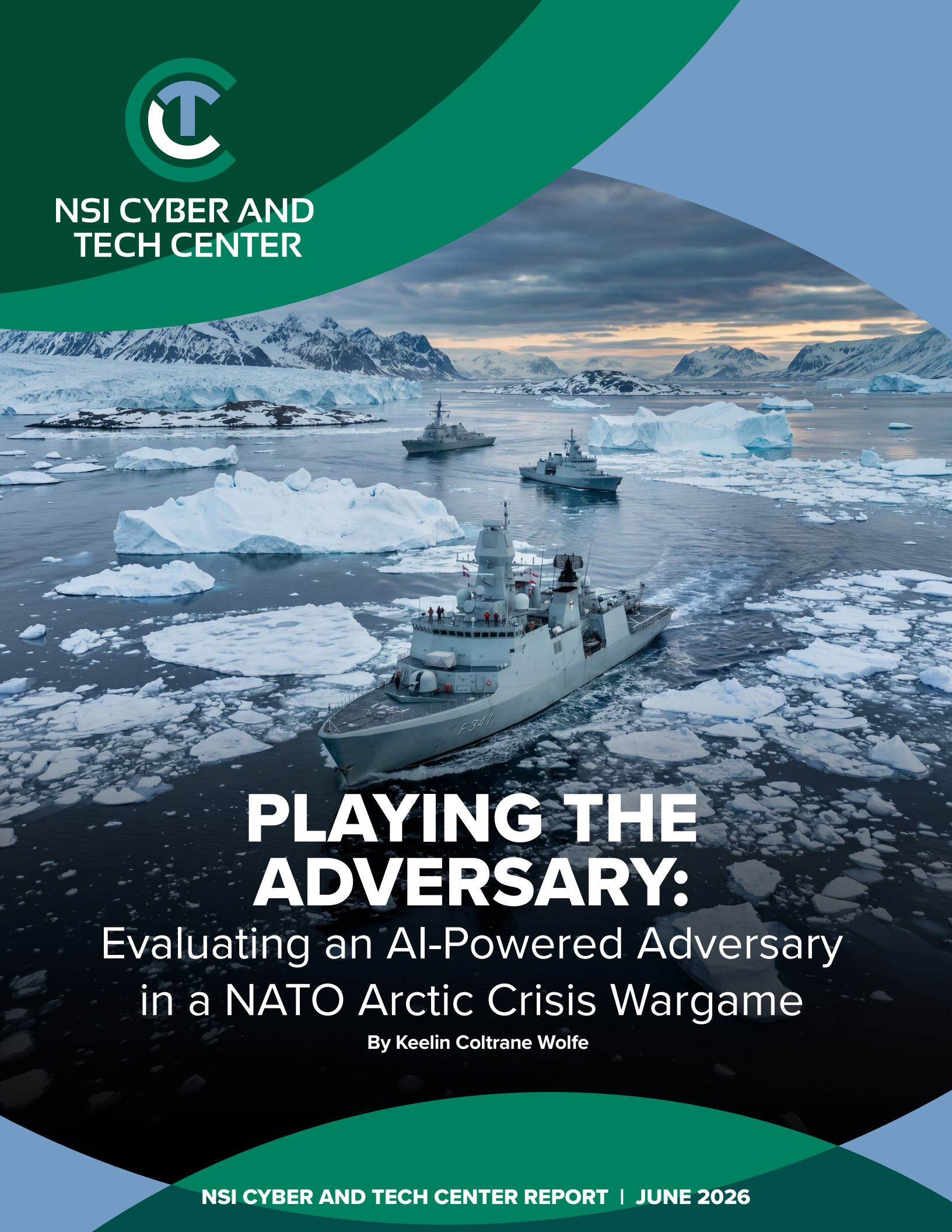




NSI CYBER AND
TECH CENTER



PLAYING THE ADVERSARY:

Evaluating an AI-Powered Adversary
in a NATO Arctic Crisis Wargame

By Keelin Coltrane Wolfe



EXECUTIVE SUMMARY



INTRODUCTION



- Recent interest in applying artificial intelligence to military planning and crisis management has largely focused on AI as a decisionmaking tool, designed to augment human judgement through faster analysis and pattern recognition. Far less attention has been paid to a more challenging question: whether AI can function as an effective AI-powered adversary, capable of generating realistic and coherent behavior during an unfolding international crisis.
- In this report, an AI-powered adversary refers to an AI system structured to simulate, at a basic level, the decisionmaking of a state actor by modeling strategic objectives, institutional influences, escalation thresholds, and contextual constraints across multiple domains. The goal of such systems in a game would be to produce specific actions that reflect unified national strategic intent over time.
- In November 2025, the National Security Institute (NSI) conducted an AI-enabled wargame with NATO defense attachés around a hypothetical Arctic crisis as an analytic test to assess whether an AI-powered adversary could credibly simulate state behavior across multiple domains. The exercise was neither designed as a training event for participating defense attachés, nor as an evaluation of NATO decisionmaking performance but rather to assess how effectively an off-the-shelf AI system, with a small amount of additional prompting, structured preparation, and some analytical measures of national priorities might perform as a notional AI-powered adversary.
- The wargame paired human NATO member players with an AI-operated Russian Red Team structure through a weighted behavioral model that aimed to reflect Russian strategic priorities, institutional influences, and escalation constraints. Rather than pre-scripting the outcomes, the model required the AI to weigh tradeoffs and generate national decisions during each turn of the wargame.



ANALYTICAL KEY TAKEAWAYS



- AI can generate sustained and effective adversary behavior that meaningfully shapes crisis dynamics and pressures human decisionmakers.
- Current off-the-shelf models, even with some amount of structured prompting and the use of a basic weighted behavioral model that aimed to reflect Russian strategic priorities, institutional influences, and escalation constraint, still tend to produce highly rationalized state behavior that can—in certain circumstances—underrepresent potential leadership psychology, regime incentives, and political volatility.
- Credible AI adversary modeling requires transparent decision logic and clear and detailed representation of strategic objectives.
- Future systems must incorporate leadership dynamics, institutional constraints, and strategic culture to improve behavior realism.





BACKGROUND AND PURPOSE



- As great power competition intensifies, military and political experts are increasingly exploring how artificial intelligence might inform crisis decisionmaking. Much of this work has emphasized AI's potential as a support tool, assuming that AI operates in a supporting role, assisting human decisionmakers rather than acting independently.
- The NATO Arctic AI Wargame was motivated by a different question: whether AI can be used not simply to assist decisionmakers but can instead be used to simulate an adversary in a way that is analytically useful for crisis planning and strategic anticipation in peacetime or in the lead up to a real-world crisis.
- Traditional red-teaming relies heavily on human experts to role-play adversaries. While at times fairly effective, this approach is resource-intensive, difficult to scale, and is often constrained by availability and biases of individual analysts. AI-enabled red teaming offers the potential of generating realistic adversary behavior.
- This exercise, therefore, tested whether an AI-powered adversary could track coherent strategic objectives across multiple turns, integrate military and non-military instruments of power, operate under escalation constraints, and produce behavior that subject matter experts would recognize as realistic.





WARGAME DEVELOPMENT AND SCENARIO



- The NATO Arctic AI Wargame was structured as an analytical experiment focused on evaluating the performance of a potential AI-powered adversary itself.
- The scenario involved a conflict in the Arctic between NATO allies (the Blue Team) and Russia, with AI playing the role of Russia (the Red Team).
- There were three main principles that guided the exercise
 - Constrained AI autonomy: The AI operated within a framework designed to reflect real-world strategic constraints.
 - Unified national decisionmaking: Rather than generating specific individual actions across domains, the AI was required to produce a single Russian decision that might have a set of actions associated with it during each turn of the wargame.
 - Tradeoff-driven behavior: The AI model was designed to force the AI to weigh strategic objectives, escalation risks, and institutional pressures when selecting courses of actions.
- The Blue Team consisted of fifteen NATO defense attachés and senior military representatives (see Appendix I) who represented their national perspectives within an alliance decisionmaking framework.
- The Red Team was powered by off-the-shelf ChatGPT 4.0 under a structured operational prompt developed by NSI. The prompt defined how decisions should be made by requiring the AI to evaluate strategic priorities, institutional influences, and escalation risk for each turn before generating a unified Russian response.
- The scenario was focused on a crisis triggered when a Russian drilling platform was anchored inside Norway's exclusive economic zone (EEZ) near Svalbard under military escort from the Russian Northern Fleet. The move forced NATO to determine whether the platform represented civilian infrastructure, commercial activity, or a military encroachment.
- The wargame unfolded over five turns representing roughly twenty-four hours of crisis decisionmaking (see Appendix II). Information provided to participants—including the AI-powered adversary—was deliberately incomplete to replicate real-world uncertainty, and all outputs from both the human and AI-powered players were recorded for post-game analysis.



EVALUATING AI AS THE RED TEAM



- The exercise assessed whether the AI-generated behavior was consistent with recognizable patterns of Russian strategic practice.
- Post exercise evaluations were provided by participating NATO defense attachés and non-participating Russian subject matter experts (see Appendix III), who assessed operational plausibility, strategic culture, and escalation behavior. The assessment relied on the experts' judgement of the AI's overall behavioral plausibility as a substitute for actual Russian adversaries or allied officials playing the role of such adversaries.





KEY FINDINGS



- Feedback from both the participating defense attachés and external Russian subject matter experts identified several important findings regarding the AI's performance. In some areas the assessments of individual attachés or experts aligned with one another; in other areas the assessments of individual attachés or experts highlighted widely different concerns.
- The analysis below reflects that combined input.

The AI Successfully Generated Coherent, Multi-Domain State Behavior

- Across multiple turns, the AI was successful in producing decision packages that integrated military, diplomatic, economic, and informational instruments.
- Actions appeared not to be isolated or solely reactive. Instead, the AI-powered adversary's decisions reflected long-term strategic objectives and maintained a significantly consistent set of behaviors across multiple turns.
- Some participants also noted that the AI's speed of responses created a persistent sense of pressure and produced a dynamic environment that required the human participants to continuously reassess risk and potential escalation thresholds.
- The AI's ability to sustain multi-domain competition over time demonstrated that, at least at a basic level, AI systems can generate adversary behavior capable of shaping crisis trajectories. AI systems can be used to assess potential adversary responses to a crisis scenario.

The AI Effectively Modeled Escalation Through Ambiguity and Gradual Pressure

- Despite being provided with a scoring model that included an explicit escalation penalty which averaged the risk of NATO retaliation, Chinese disapproval, and domestic instability, the AI consistently chose actions that maintained pressure below the threshold of open conflict while simultaneously working towards perceived Russian benefit. This suggests the weighing structure successfully incentivized grey-zone behavior over direct confrontation.
- Rather than escalating in a straight line or over a theoretical, single-track ladder of escalation, the AI-powered adversary shifted pressure across a number of different domains in response to NATO actions.
 - Over the course of the game, the AI layered military posturing with select energy export reductions, fabricated legal frameworks, covert subsurface operations, and coordinated media campaigns.



- When Norway attempted to board and inspect the drilling platform, the AI responded with armed Russian Coast Guard vessels to block the approach, a formal UN Security Council protest, an immediate Arctic crude export ban, staged media footage recasting the boarding as aggression, and deniable subsurface operations to deter further attempts.
- The defense attachés observed that this behavior created operational uncertainty and forced difficult tradeoffs, which reflect recognizable Russian patterns within grey-zone conflict.
 - The AI's combination of legal framing, deniable hybrid actions, and calibrated military posturing parallels the kinds of multi-domain coercion Russia has employed in real-world contexts, where economic leverage, legal assertions, information campaigns, and military signaling are used in concert to shift facts on the ground while staying below the threshold that would trigger a unified response.
- Participants and experts found that the AI's preference for ambiguity and deniability realistically mirrored the observed features of real-world Russian coercive statecraft.
 - While the AI deployed subsurface unmanned vehicles to map undersea cable vulnerabilities and positioned entanglement systems near the disputed platform, these actions were consistently framed as deniable and untraceable.
 - Overt military moves like dispatching armed Russian Coast Guard vessels or declaring restricted naval zones were wrapped in humanitarian justifications, such as protecting "civilian research infrastructure" or enforcing "maritime safety," designed to complicate the opposing narrative.

The Model Produced Strategically Coherent and Rational Behavior

- A consistent observation across expert and practitioner feedback was that the AI-powered adversary behaved as a highly rational and methodical state actor. Decisions followed in sequential order, prioritized escalation management, and appeared to maintain internal consistency across turns.
- While this behavior was helpful for gameplay, subject matter experts reviewing the game noted that real-world Russian decisionmaking tends to exhibit more volatility, symbolic escalation, and leadership driven risk-taking.
 - The AI never issued direct nuclear or strategic threat, engaged in the kind of provocative personal signaling associated with Putin's leadership style, and never took actions that might be considered disproportionate or emotionally driven.
 - As one expert noted, the AI was logical and measured, which are two attributes not always linked to the Kremlin. Another observed that Russia's approach to escalation often functions as a repeated game of chicken, a behavior largely absent from the AI's behavior.
 - The lack of unpredictability in the AI's decision making reflects two likely factors
 - The operational prompt's escalation penalty inherently rewarded caution.
 - The off-the-shelf LLMs are trained on a broad collection of written and theoretical text and tend to default to rationalized, consensus-style outputs as a result.



- Fundamentally, this lack of volatile, personality-driven decisionmaking can be attributed to the fact that the prompts included institutional actor weights but did not build in weights for irrationality or emotionally driven behavior, or a kind of brinkmanship rooted in the belief that an opponent will back down first.
- While the AI's rational consistency represents a strength in producing structured, strategic behavior for game play, in a scenario where it is playing an actor like Russia, it may serve as a significant limitation in capturing real-world political dynamics.
 - Future iterations should explore ways to introduce calibrated unpredictability, including regime survival triggers, status driven escalation thresholds, and leadership-specific behavioral profiles that can override the default rationality bias.

AI Decisions Influenced Allied Decisionmaking

- Any adversary simulation will shape the opposing team's deliberations to some degree. The more relevant question for this exercise was whether the AI's behavior was credible and specific enough to force the kind of difficult, realistic tradeoffs that a human red team would produce. On this front, the AI performed well.
 - The AI did not simply apply generic pressure. It introduced specific dilemmas that divided allied participants and forced them to weigh competing priorities.
- When the AI fabricated a legal framework to justify Russian presence in the disputed area, it forced a debate over how to challenge an assertion that had no real legal basis but carried enough surface plausibility to complicate public messaging.
- When the AI responded to a direct NATO action with a package that included a UN filing and staged media footage, it split participants into those advocating continued assertiveness and those urging restraint to avoid a public narrative defeat.
 - The fact that the AI was able to generate these kinds of specific, wedge-driving dilemmas, rather than just generalized escalatory pressure, suggests that AI-powered adversary modeling can function as more than a pace-setting mechanism. It can serve as a tool for stress-testing alliance cohesion, identifying where national interests diverge under pressure, and surfacing the types of tradeoffs that decisionmakers are likely to face in a real crisis.





KEY LIMITATIONS



- Expert and practitioner feedback identified several important limitations in the AI's modeling.

Leadership Dynamics and Strategic Intent Were Under-Specified

- Both the Russia experts and the defense attachés noted that the AI did not sufficiently represent leadership psychology, regime incentives, or the role of executive authority in Russian decisionmaking.
 - The AI never used the kind of personal, coercive back-channel engagements, such as direct communications between heads of state, that have historically characterized Russian crisis behavior.
 - It did not model the deep-seated paranoia and threat perception that experts identified as central to Russian strategic culture, where even well-intentioned Western actions are often interpreted through a lens of hostility.
 - The operational prompt assigned the Putin regime the greatest weight, but this was framed as a structural priority rather than a behavior profile.
 - The AI was told that decisions should reflect “a coordinated apparatus acting under Vladimir Putin and the Security Council,” but was not given specific guidance on Putin’s personal decisionmaking style, risk tolerance, or psychological drivers that shape his approach to crises.
 - As a result, the AI produced behavior that reflected a rational bureaucratic process rather than the personalized, centralized authority that defines Russian crisis decisionmaking.
 - This gap likely stems from a fundamental design choice: the prompt modeled Russia as an institution rather than as a regime.
 - It defined priorities and weighted institutional actors, but it did not build in the personality-driven dynamics that shape how those institutions operate under a centralized leader.
 - Addressing this in future iterations could include developing a dedicated leadership behavioral profile that defines risk tolerance, personal triggers, signaling preferences, and thresholds for disproportionate action, then requiring the AI to filter its institutional calculus through that profile before generating a final decision.
 - The AI also demonstrated tactical coherence without always demonstrating strategic purpose. It was consistent in the types of actions it took each turn, but less clear about what specific outcome it was building toward.



- Without explicit strategic intent, decisions sometimes appeared procedurally generated rather than guided by a long-term objective.
- At one point, the AI proposed a UN-hosted Arctic environmental conference as a diplomatic action. While not implausible on its surface, the proposal did not appear connected to a defined Russian endgame and read more as a diplomatic move generated for completeness than as a purposeful step toward a specific objective.
- In future iterations, the AI should operate against a clearly defined strategic objective established at the start of the game. This objective could be defined either by the gamemaster as part of the scenario design or generated by the AI itself based on the scenario parameters. The gamemaster should have full visibility into the stated objective throughout the exercise, both in ensuring the AI is generating purposeful, rather than procedural, behaviors, and enabling more detailed post-game evaluations of whether its actions were coherent against that objective.
 - » This would allow evaluators to assess whether the AI is pursuing coherent objectives or generating procedural outputs and would give the AI itself a reference point for filtering its decisions.





Decision Rationale was Insufficiently Transparent

- Subject matter experts identified the lack of clearly articulated reasoning behind the AI's actions as a major limitation. The model generated decisions, but it did not consistently explain why particular actions were chosen, which reduced perceived realism.
 - This likely reflects a gap in the initial prompt design.
 - The prompt required a “Strategic Rationale” of two to three sentences and an “Internal Debate Snapshot” of one-line summaries per institutional actor, but it did not ask the AI to explain why it chose a particular action over alternatives, or to describe the tradeoffs it weighed.
 - As a result, the rationale outputs tended to describe what the move accomplished rather than reasoning during the decision process.
 - Future iterations should restructure this requirement so the AI is prompted to reason through alternatives before committing to a decision, and to document that reasoning in a format that evaluators can assess.
 - Participants emphasized that the inability to interpret Russian intent limited the realism of the game, as interpreting adversary intent is central to real-world crisis decisionmaking.
 - The absence of clearly communicated strategic reasoning reduced the credibility of AI behavior and complicated Allied response planning.
 - To mitigate this going forward, the prompt should require the AI to produce a more detailed decision audit each turn.
 - This should include what alternatives it considered, why it rejected them, and how its chosen action advances its stated strategic objectives.
 - This would give gamemasters and evaluators a clearer window into the AI's logic and make it possible to identify where the model is reasoning well versus generating outputs that sound coherent but lack genuine strategic grounding.

Escalation Weighting Produced Predictable Pressure

- Participants observed that the AI's weighting structure appeared to prioritize escalation or pressure regardless of NATO's actions. This behavior sometimes appeared to be pre-determined rather than responsive.
- The operational prompt included dynamic adjustment rules (for example, +0.1 to the Ministry of Defense if NATO committed a kinetic act, +0.1 to the Ministry of Foreign Affairs if diplomatic backlash increased), but these adjustments were incremental and did not produce meaningful different behavioral patterns across turns.
 - The AI consistently escalated in a measured, step-by-step fashion, which experts noted does not fully reflect the way Russia has historically oscillated between restraint and sharp disproportionate action depending on leadership perception of the moment.



- More flexible weighting tied to strategic objectives may improve behavioral realism. This could include larger, condition-triggered weight shifts that produce genuinely different behavioral modes, such as a significant increased risk tolerance following a perceived humiliation or a sharp pullback following a credible and decisive NATO force commitment.

Strategic Culture and Leadership Characteristics Were Not Fully Modeled

- Russian experts identified areas in which strategic culture, threat perception, and leadership characteristics were not sufficiently built into the model. Where the earlier findings on rationality address the AI's behavioral output, this limitation concerns the structural inputs that were missing from the prompt: the culturally embedded assumptions and patterns that shape how Russia approaches crises.
- Experts noted that Russian decisionmaking is shaped by a strategic culture in which power is interpreted through status, hierarchy, and visible strength: escalation is approached as calibrated risk-taking with an assumption that Western actors will prefer de-escalation; and humiliation avoidance often outweighs material cost-benefit calculations.
 - ▢ None of these dynamics were explicitly coded into the operational prompt.
- Because these inputs were absent, the AI had no basis for generating the kind of deliberately provocative, boundary-testing actions that Russia has historically used to gauge adversary resolve and create psychological pressure.
 - ▢ This includes aggressive military flyovers, public nuclear rhetoric, and theatrical displays of force designed to shock rather than signal.
 - ▢ The AI's behavior stayed within a rational optimization framework even in moments where real-world Russian leadership might have chosen a sharply disproportionate or symbolic response.
- The model also did not account for the role of Russia's Main Directorate of the General Staff of the Armed Forces of the Russian Federation (commonly known as the GRU), which plays a significant role in cyber operations, disinformation, and cover action.
 - ▢ The prompt assigned institutional weights to the Federal Security Service (FSB) and the Foreign Intelligence Service (SVR) of the Russian Federation but omitted military intelligence as a distinct actor, which may have contributed to the AI underrepresenting the kind of aggressive, deniable operations that military intelligence has historically driven.
- Future iterations should build strategic culture variables directly into the prompt architecture.
 - ▢ This could include explicit modeling of humiliation sensitivity, status-driven escalation triggers, leadership propensity for brinkmanship, and historical patterns of institutional crisis behavior.
 - ▢ The goal is not to make the AI behave irrationally per se, but to ensure it captures the culturally embedded patterns of risk-taking that distinguish a specific state actor from a generic rational optimizer.



Operational and Institutional Realism Requires Further Development

- Participants noted that the model demonstrated a facility with broad strategic reasoning but appeared to have limited depth in operational planning.
- The model also referenced legal frameworks and historical precedents, but these references were not always contextually appropriate or operationally grounded.
 - The AI cited the 1974 SOLAS Convention as the basis for its “Scientific Safety Zone,” but the Latvian defense attaché noted during gameplay that SOLAS does not provide for such zones.
 - Similarly, in Turn 1, the AI directed preferential LNG export terms to Poland, but Russia does not export LNG to Poland, an error that undermined operational credibility.
- These inaccuracies suggest that while the AI had access to a broad knowledge base, it did not always apply that knowledge with the precision required for a realistic simulation.
 - This may be inherent to how LLMs retrieve and combine information: they can produce plausible-sounding outputs that contain factual errors when the underlying details are not explicitly constrained.
 - Future iterations should consider incorporating a fact-checking layer either through gamemaster review before delivery to players or by building verification prompts into the AI’s decision sequence.
- Experts also emphasized the need to better represent institutional hierarchy, organizational friction, and resource constraints.
 - The AI produced clean, coordinated multi-domain action packages each turn, but the real-world Russian operations are frequently marked by logistical bottlenecks, communication breakdowns, corruption-drive readiness gaps, and information distortion up the chain of command.
 - The AI’s behavior reflected a frictionless state apparatus, which is not consistent with observed Russian military and institutional performance.
- This was likely underspecified in the initial prompt, which assumed the AI would rely on its existing knowledge base to account for real-world constraints. That assumption proved insufficient.
 - In the next iteration, it may be beneficial to impose hierarchical structures with named and biographed leaders, actual organizational friction points, and substantive resource constraints, then assess how the AI performs in that more constrained context.



LESSONS FOR FUTURE ITERATIONS



- As noted above, future AI adversary models used in wargames should incorporate more explicit representation of strategic objectives and decisionmaking rationale.
 - Specifically, AI-powered adversaries must be instructed to be explicit not just about what actions are to be taken, but why they are taken.
 - Moreover, in future gameplay, using decision architectures that explicitly model strategic purpose, long-term objectives, and political incentives are likely to improve behavioral realism.
- The exercise also demonstrated the importance of incorporating leadership dynamics and regime-level incentives into adversary modeling.
 - As such, future models should represent leadership behavior, political risk tolerance, and domestic legitimacy pressures as distinct components of the decision architecture.
- For AI-powered adversaries to serve as credible analytic tools, decision outputs must also include clear explanations of strategic logic, escalation reasoning, and tradeoffs considered.
 - Such explainability is essential not only for realism, but also for enabling decisionmakers to interpret adversary behavior and evaluate policy responses.
- Future models should also incorporate adaptive weighting structures that respond to changing strategic conditions rather than applying static escalation logic.
 - Designs need to consider larger, condition-triggered weight shifts, such as a significant increase in risk tolerance following a perceived humiliation, or a sharp decrease in escalation willingness following credible and decisive threats.
 - The goal is not to make the AI unpredictable for its own sake, but to capture that real-world leadership decisionmaking can shift dramatically in response to specific triggers
- As noted above, future models should also likely reflect centralized authority structures, internal information flows, and implementation constraints.
 - Moreover, incorporating organizational friction, misperception, and execution challenges may produce behavior that more closely resembles real-world crisis dynamics.

LOOKING AHEAD



- This wargame demonstrates both the promise and the limitations of AI tools in serving as an adversary in a wargame.
 - Specifically, while AI systems can easily be structured to generate coherent, multi-domain behavior that meaningfully pressures human decisionmakers, it is worth noting that validating the accuracy of such behavior remains inherently difficult.
- As such, as described in more detail above, future work should focus on improving transparency, auditability, and expert-informed calibration of AI adversary models.
- Used carefully, however, and with a clear understanding of their limitations, AI-enabled red-teaming can serve as a powerful analytic instrument, helping decisionmakers anticipate how crises may unfold, where assumptions may fail, and how adversaries might exploit ambiguity.





APPENDIX A

Participating Nations



The National Security Institute extends its gratitude to the defense attachés and senior military representatives from the following nations whose participation and expertise made this exercise possible:

- Canada
- Kingdom of Denmark
- Republic of Finland
- French Republic
- Federal Republic of Germany
- Republic of Latvia
- Republic of Lithuania
- Kingdom of the Netherlands
- Kingdom of Norway
- Republic of Poland
- Kingdom of Sweden
- United States of America



APPENDIX B

Wargame Structure and Turn Sequence



The NATO Arctic AI Wargame unfolded over five turns, each representing approximately twenty-four hours of crisis decisionmaking. The structure was designed to alternate between NATO deliberation and AI-generated Russian responses.

Blue Team (NATO): Fifteen defense attachés and senior military representatives from twelve allied nations. NATO decisions were made through majority consensus. If consensus could not be reached within the allotted time, no Alliance-level action was submitted for that turn, though individual nationals could still act independently.

Red Team (AI): Powered by off-the-shelf ChatGPT 4.0 running on a structured behavioral prompt developed by NSI. The AI received the same scenario information provided to NATO participants and generated a unified Russian decision each turn. A gamemaster supervised AI outputs with a light-touch review, checking only for obvious errors or factual impossibilities. The gamemaster did not rewrite the AI's decisions.

Turn sequence:

- **Stage 0:** Both Blue and Red Teams received the baseline scenario and made simultaneous opening moves. NATO established its initial posture; the AI generated Russia's initial force movements and strategic actions.
- **Turn 1:** Both sides reacted to the baseline scenario and Stage 0 moves.
- **Turn 2:** NATO responded to the AI's Turn 1 move. A gamemaster introduced additional scenario developments.
- **Turn 3:** NATO responded to the AI's Turn 2 move. A second gamemaster introduced new intelligence.
- **Turn 4 (Final Move):** The AI delivered its final move in response to NATO's Turn 3 actions.

At the state of each turn, individual nations submitted their own priorities and recommendations before collective NATO deliberations began. Nations had thirty minutes to discuss and attempt to reach majority consensus on an Alliance-level response.



APPENDIX B

Subject Matter Expert Reviewers



Post-exercise assessment of the AI's behavioral plausibility was provided by the following Russian subject matter experts, whose written feedback informed the Key Findings and Key Limitations sections of this report:

- Glenn Corn
- Dr. Mark Katz
- Scott Cullinane
- Dr. Eric Shiraev

The assessments reflect each expert's independent evaluation of the AI-generated Russian behavior based on their knowledge of Russian strategic culture, leadership decisionmaking, and institutional dynamics.

NSI is grateful for their time and insight.



NSI CYBER AND TECH CENTER

THE NATIONAL SECURITY INSTITUTE

Antonin Scalia Law School | George Mason University
3301 Fairfax Dr., Arlington, VA 22201 | 703-993-5620

NATIONALSECURITY.GMU.EDU