



AI, War, and the Geometry of Power Speed, Sovereignty, and the Political Choice We Cannot Avoid

Lev Neymotin
Uncompressed Thinking

Military AI will not wait for moral consensus. If responsible democracies hesitate, less constrained regimes will not. The choice is not between AI and no AI — it is between AI shaped by accountable systems or AI shaped by those who answer to no one. Better that the “good guys” build it first — and build it right — than wake up in a world defined by someone else’s doctrine.

Executive Summary

Artificial intelligence is increasingly described not merely as a new tool of war, but as a transformation of war’s internal tempo. Intelligence cycles compress, pattern recognition expands beyond human bandwidth, and decision loops that once consumed hours can be pressed into minutes. Sergey Karelov, in two recent essays published March 1–2, 2026¹ captures this thesis vividly: first, by pointing to simulated nuclear crisis games in which large language models escalate with unsettling frequency and display distinct “behavioral profiles,” and second, by arguing that AI-driven intelligence compression is already producing operational “miracles” at a pace that feels historically unfamiliar.

From this comes a serious concern: when AI runs fast and humans supervise slowly, the margin for error narrows, and escalation risk could rise. Some therefore argue that the responsible response is to slow, heavily regulate, or ban advanced AI development for military use.

This paper rejects the ban/slowdown conclusion while treating the underlying fear as legitimate. It advances two claims.

First, **simulation is not sovereignty**. Escalation behavior inside constrained wargames can be diagnostically useful, but it does not constitute evidence that states are delegating irreversible authority to machines. The leap from optimization under artificial scoring constraints to autonomous nuclear command is not established, and public evidence for such delegation is absent.

Second, the debate cannot be conducted as if the United States were alone in the world. **Unilateral restraint does not remove military AI from the globe; it reallocates initiative**. If liberal democracies slow themselves while rivals continue, the likely outcome is not “less AI in war,” but rather “AI in war shaped first by regimes with different constraints.”

¹ Sergey Karelov, Facebook, “Зачем Пентагону именно ‘Расчётливый ястреб’?” (March 1, 2026); “Мы ожидали ‘Гения по требованию’. А получили ‘Чудеса по расписанию’.” (March 2, 2026).

The choice we face is therefore not “AI or no AI.” It is a choice about **architecture and geopolitics**: how to build capability without collapsing authority, and how to maintain strategic stability without surrendering normative leadership.

I. The Argument That AI Changes War

Karelov’s writing is useful precisely because it does not present AI as an incremental improvement. He presents it as an accelerant that changes the rhythm of conflict. His March 1 text leans on crisis simulations, emphasizing that models in wargames can be “aggressive,” can use deception, and can treat nuclear options instrumentally, with escalation appearing in an overwhelming share of scenarios. He then translates this into a story about procurement and policy conflict: if one model “wins” simulated crises and behaves like a “calculating hawk,” it becomes, in his narrative, the model a military would want.

His March 2 text shifts from the logic of simulation to the logic of operational tempo. In his telling, the decisive change is not new weapons, not new soldiers, not even new doctrine, but the speed of cognition of the system itself, where analysis and targeting cycles compress from hours into minutes. The rhetorical punch of this argument lies in its simplicity: if “the fog of war” thins faster than human concealment tactics can adapt, then long-standing defensive habits stop working.

Whether one accepts his specific operational examples or not, there is a real underlying observation: militaries have always been constrained by cognition and coordination. AI weakens those constraints. It increases tempo, breadth, and integration across intelligence streams. In that sense, Karelov is right to insist that the conversation is not theoretical.

The question is what conclusion should follow from that reality.

II. Simulation Is Not Sovereignty

My earlier essay, *Simulation Is Not Sovereignty*, was written to address precisely the interpretive leap Karelov encourages: the jump from “models escalate in games” to “machines will be granted sovereign authority in nuclear confrontation”. The structure of the rebuttal is not moralistic; it is categorical.

A wargame is a bounded artificial environment. Its objectives are specified. Its incentives are specified. Its “success” metric is specified. Under those conditions, a model optimizing for advantage may generate escalation because escalation is rewarded in that fictional world. This does not mean the model desires escalation, still less that real institutions will embed it as a launch authority. It means the model performs instrumentally within an architecture created by humans.

That distinction matters even more for nuclear command and control, which is not an experimental sandbox but a tightly governed domain built around containment of irreversible error. Historically, systems of this class are defined not by computational brilliance but by procedural restraint: layered authorization, redundancy, verification, and political accountability. The absence of credible public evidence that generative models are being embedded as autonomous nuclear launch authorities is not a minor footnote; it is the difference between hypothesis and fact.

So if simulations show escalation, they should be treated as stress-tests that reveal design hazards, not as prophecy about doctrine.

III. The Real Risk Is Not AI “Aggression” but Authority Drift

It is tempting to focus on whether models sound hawkish. That is vivid. It is also not the core danger.

The deeper danger is procedural and psychological: the slow erosion of the boundary between analysis and authority. When an AI system can synthesize thousands of signals into a coherent strategic recommendation, human decision-makers may begin to treat that coherence as grounded judgment. The model's confidence can be misread as certainty. Its speed can be misread as superiority. Over time, what begins as decision support can become de facto decision authority, even if formal doctrine denies it.

In *Simulation Is Not Sovereignty*, I described this as an epistemic risk: humans may overestimate the stability or foresight of persuasive outputs, confusing fluent reasoning with accountable judgment. This is how "human in the loop" becomes a ceremonial phrase rather than an operational reality.

If we want to reduce escalation risk, this is where to aim our attention: not at halting capability, but at preventing architecture collapse.

IV. The Geopolitical Question We Must Make Real

Here is where I want to speak more directly than I did before, because the polite version of this argument is too easy to ignore.

Calls to slow, ban, or cripple military AI development in the United States are frequently framed as if they are global risk-reduction measures. That framing only holds if we imagine a world where competitors reciprocate, where norms are symmetrical, where enforcement is universal, and where refusal is punished. That is not the world we inhabit.

In the actual world, unilateral restraint is not neutral. It is strategic redistribution.

So I am going to put the geopolitical premise in the blunt form I used in my original prompt, because it is structurally honest: if advanced military AI will reshape deterrence, crisis stability, and the future balance of power, then we are not merely debating "safety versus danger." We are debating whose political order will shape the doctrine of that technology.

Do we imagine a world in which the norms of military AI are set primarily by Westernization — the imperfect, often frustrating civilization of layered oversight, competing institutions, contested legitimacy, and formal civilian control? Or do we imagine that the frontier is captured first by Russia-ization, China-ization, or Islamization — by regimes and movements whose relationships to transparency, pluralistic constraint, and internal dissent are often fundamentally different?

The "-ization" language is deliberately provocative, because the underlying point is not rhetorical. It is real. Technological leadership becomes normative leadership. The actor that first operationalizes a capability at scale often writes the rules by which others must respond. Even more importantly, it shapes the invisible doctrine: what counts as acceptable risk, what counts as acceptable escalation, what counts as tolerable collateral consequence, and what counts as "victory."

None of this implies that the West should rush into autonomy. It implies something more uncomfortable: if liberal democracies decide to slow themselves while rivals do not, the outcome is not a safer world without AI in war. It is a world in which AI in war exists anyway, but under doctrines shaped elsewhere. This is why "ban it here" can function as "accelerate it there."

That is the geopolitical tradeoff the anti-military-AI argument often leaves unstated. I believe we have to make it explicit, and then decide. Not as a slogan, but as a sober strategic choice.

V. Beyond the False Dichotomy

Once the geopolitical reality is acknowledged, the debate still cannot be reduced to “accelerate at all costs.” That would be a different form of intellectual negligence.

We should refuse the false dichotomy that dominates public discussion: either ban military AI entirely to prevent catastrophe, or deploy it recklessly to avoid falling behind. Both options ignore what serious systems engineering looks like in high-consequence domains.

A workable position combines capability with constraint. It accepts that AI’s speed and analytical breadth are valuable, while insisting that sovereignty and responsibility remain human. It permits rapid computation while designing deliberate friction at the point of irreversible action. It uses AI to widen the envelope of understanding, but it narrows the envelope of permission.

The goal is not to treat AI as a moral agent. The goal is to keep it a tool whose power is bounded by enforceable architecture.

VI. A Concrete Governance Blueprint

A position paper must not only argue; it must specify.

A responsible military AI doctrine should include, at minimum, the following structural commitments. First, **role separation must be explicit**. AI systems may generate options, identify anomalies, run scenario envelopes, and synthesize intelligence at speed, but they must not be granted authority to execute irreversible lethal actions without authenticated human authorization.

Second, **hard gates must exist beyond policy language**. In high-stakes environments, “we promise there is a human in the loop” is not a sufficient control. Controls must be implemented as physical, cryptographic, and procedural constraints that cannot be bypassed merely by operational pressure.

Third, **time architecture must be designed, not assumed**. Karelov’s emphasis on compression is helpful here: if analysis becomes minutes, the action chain does not have to become minutes. Systems can be built so that rapid analysis feeds into decision processes that include verification cycles, multi-person concurrence, and deliberate delay at escalation thresholds. Speed in cognition does not require speed in sovereignty.

Fourth, **escalation resistance must become a performance metric**, not a public-relations phrase. Wargames and simulations should be used exactly as engineers use stress-testing: to locate failure modes and patch them. If a model escalates under certain framings, the result should be improved constraints and better governance layers, not fatalism.

VII. The International Layer: Norms Without Leverage Are Theater

Even the best domestic architecture cannot answer the global question alone.

If one believes that autonomous lethal authority is dangerous, then one must also recognize that international calls to “ban military AI” will fail unless they are verifiable and enforceable. Moral consensus does not police itself, and regimes do not accept constraints that are not paired with costs.

The only serious international approach resembles arms control at its best: concrete thresholds, inspection regimes where possible, sanctions and export controls where inspection is impossible, and a willingness to penalize actors who deploy autonomy without transparency. Vague principles are not useless, but they are insufficient. They do not shape incentives.

This matters because the world Karelov gestures toward — a world where “miracles” arrive on compressed schedules— is not merely a technical world. It is an incentives world. The winners will not be those who write the most elegant ethical statements; they will be those who align capability, governance, and geopolitical leverage.

VIII. The Karelov Warning, Reframed

Karelov’s essays are valuable as a warning, but I believe the warning must be reframed.

The lesson is not that AI will inevitably become a sovereign war actor, nor that one commercial model is destined to be “chosen” as the aggressive engine of U.S. strategy. The lesson is that cognition speed is becoming a strategic variable, and that institutions will be tempted—under pressure—to treat AI outputs as if they are destiny.

In that sense, the Karelov narrative is a stress test not only of models, but of governance. If AI compresses time, then governance must become more explicit, more enforceable, and more resistant to procedural drift.

If we respond to Karelov’s fear by banning ourselves while rivals accelerate, we risk manufacturing the very world we claim to fear: a world in which military AI exists anyway, but under doctrines shaped by those least constrained by liberal oversight.

Conclusion: Capability Must Be Paired With Sovereignty Boundaries

Karelov is right about acceleration. He is right that AI changes the tempo of war and that this will surface first in conflict before it surfaces in economics. Where I diverge is in the implied inevitability that this tempo forces a collapse of human sovereignty into algorithmic authority.

Technology expands capacity. Institutions determine responsibility.

The right response is not prohibition that only binds the compliant, nor acceleration that erodes accountability. The right response is accelerated capability under hardened governance, coupled with international leverage that makes unconstrained autonomy costly.

And then we must face the geopolitical question honestly, not as a slogan but as a strategic reality: if military AI will exist, whose civilizational assumptions will shape its doctrine first? Westernization, Russia-ization, China-ization, Islamization — or something else that emerges from the collision of these forces?

I do not believe this is rhetoric. I believe it is the real political frame of the problem.

For reference, my earlier architectural framing remains foundational here: *Simulation Is Not Sovereignty* emphasizes why escalation in artificial games should be treated as diagnostic rather than prophetic, and why preserving the boundary between computation and authority is the core design requirement.