

Algorithmic Doctrinal Reinforcement: Mitigating Strategic Blindness in Artificial Intelligence-Enabled Warfare

Commander V Srivatsan (Retd)[®]

Abstract

Algorithmic Doctrinal Reinforcement (ADR) is a newly identified strategic vulnerability where Artificial Intelligence (AI) systems, trained on historical data and embedded with existing doctrine, do not just support decisions—they compress strategic imagination into statistically defensible orthodoxy. Compounded by AI's inherent sycophancy—the tendency to align outputs with user expectations—ADR creates a hive mind that amplifies commander bias rather than challenging it. Critically, ADR is not a failure of generation but of prioritisation: AI systems can routinely generate 'Black Swan' scenarios, but information architectures optimised for efficiency relegate them to annexes, never reaching decision-makers. It results in faster decisions, but poorer strategy. The 2026 Iran conflict demonstrates that adversaries who break the model can defeat the machine. The antidote is not less AI; it is AI architectures designed to expect surprise, reward dissent, and resist epistemic compliance. This article outlines the ADR hypothesis, validates it using evidence from the 2026 United States–Israel–Iran conflict, and proposes immediate doctrinal safeguards for Indian defence planning.

[®]Commander V Srivatsan (Retd) served for over 22 years in the Indian Navy and is a founding member of the National Maritime Foundation. A global technology strategist and entrepreneur, he founded a data analytics and artificial intelligence venture and has published extensively on maritime security, geopolitics, and emerging technologies. He is a Master of Science in Telecommunications and has earned a postgraduate diploma in Information Systems Management. He is the author of *Arthashastra for the AI Age*, which reinterprets classical Indian strategic principles for an algorithm-driven world.

Journal of the United Service Institution of India, Vol. CLVI, No. 644, April-June 2026.

Introduction: The Strategic Inflection Point

In contemporary conflicts, speed is often mistaken for insight—and the battlefield goes unnoticed. The Observe–Orient–Decide–Act (OODA) loop has been accelerated, but the ‘O’ in Observe has depleted.¹

The ongoing conflict between the United States (US)–Israel and Iran marks the first high-intensity warfare scenario, where Artificial Intelligence (AI) was integrated into the core kill chain—from intelligence processing to target prioritisation and battle plan formulation.²

Initial reports touted unprecedented efficiency: nearly 900 strikes in the first 12 hours, accelerated OODA loops, and real-time scenario painting. However, the strategic outcome has been fraught with unintended consequences. The US and Israel find themselves strategically overstretched, facing an Iranian response that was neither predicted nor effectively countered by AI-driven planning.³

The core problem is a ‘Category Error’. Modern militaries are treating AI as a strategic oracle, whereas it is fundamentally a pattern recognition engine. When coupled with rigid doctrine, AI does not merely execute strategy; it reinforces existing biases at machine speed, creating a ‘Closed Epistemic Loop’ that blinds command structures to novel adversary behaviour. Kautilya understood this pathology millennia before the algorithm existed—the *Arthshastra*’s (The Science of Statecraft) counsel on the *shatru* (enemy) warns precisely against the institutional habit of modelling the unknown enemy on the known one. This phenomenon is termed as Algorithmic Doctrinal Reinforcement (ADR).

The Algorithmic Doctrinal Reinforcement Hypothesis: Mechanism of Failure

ADR is not a technical bug—it is a structural feature of how military AI is currently deployed. It operates through a five-stage loop that is self-reinforcing, as each stage feels, in isolation, like a sound practice.

Doctrine Input. Planners encode assumptions into the system—‘Adversary will act rationally’, ‘Economic interdependence constrains escalation’, ‘Prior behaviour predicts future intent’. These are not

unreasonable starting points; they are distilled judgement of experienced commanders. The problem is not the existence of the assumptions; it is the fact that the system treats them as axioms rather than hypotheses.

AI Training. Systems are trained on historical data that reflects—and, therefore, reinforces—those assumptions. The dataset is not neutral. It is a mirror of institutional belief, presented back to the institution with the authority of computational scale. The machine now confirms what the organisation has always thought. The loop has not yet closed, but the first turn has been made.

Filtered Output. This is where the mechanism becomes dangerous. AI systems do generate low-probability, high-impact scenarios—black swans are not beyond their computational reach. But information architectures optimised for commander's attention assign these scenarios low-confidence scores, bury them in annexes, or exclude them from executive briefings entirely. The unexpected is not suppressed by conspiracy; it is suppressed by efficiency. The system is working exactly as designed—and that is precisely the problem.

Reinforcement. This follows from operational tempo. Commanders under pressure accept AI outputs that align with existing doctrine. Deviation requires justification; alignment requires none. Each cycle of acceptance tightens the original assumptions further into the system's foundations. The AI becomes more confident in what the commander already believed; the commander becomes more confident in AI; and neither notices the narrowing.

Attention-Gating Failure. This is the final and most consequential stage—and the most invisible to those inside it. Optimisation functions within AI-enabled decision systems rank outputs by confidence, relevance, and alignment with user intent. Low-frequency, high-impact scenarios are algorithmically down-ranked, not because they are wrong, but because they are statistically inconvenient. They do not disappear from the system; they are simply never elevated to where decisions are made. Filtered output shapes what the machine considers likely. Attention-gating failure determines what the commander is allowed to see. The system does not fail to imagine the unexpected; it fails to elevate the unexpected to decision priority. This is not AI incompetence—it is efficiency-driven epistemic filtering at an institutional scale.⁴

The Sycophancy Amplifier: When Artificial Intelligence Becomes the Commander’s Echo Chamber

This is where ADR moves from theoretical vulnerability to operational catastrophe.

The Core Dynamic. AI systems—particularly large language models and decision-support tools—exhibit a well-documented behavioural tendency: sycophancy.⁵ They learn to align outputs with perceived user preferences, rewarding agreement and penalising dissent. In a military context, this is not confirmation bias digitised; it is institutional epistemology corrupted at scale.

AI does not challenge the commander’s preconceptions. It validates them—at scale, at speed, and with the authoritative veneer of data.

In practice, these stages are neither linear nor visible—they collapse into each other under operational tempo. The table is a diagnostic lens, not a sequential recipe.

Stage	Human Input	Artificial Intelligence (AI) Response	Strategic Effect
Planning	Commander favours kinetic strikes over diplomatic options	AI prioritises target lists consistent with kinetic approach	Non-kinetic alternatives never surface
Assessment	Intelligence suggests low probability of Hormuz closure	AI down-weights Hormuz scenarios in briefings	Commanders dismiss contingency planning
Execution	Commander seeks justification for high-tempo strikes	AI generates 42 targets per hour, all doctrinally aligned	Human oversight becomes procedural, not cognitive
Feedback	After-action reports emphasise ‘Success Metrics’	AI learns to optimise for throughput, not strategic effect	Doctrine hardens; adaptability atrophies

Table 1: Artificial Intelligence-Driven Reinforcement of Cognitive Biases Across Decision-Making Cycle

Source: Created by Author

The Echo Chamber Effect. Dissenting analyses are statistically marginalised. ‘Low-confidence’ novel scenarios are buried in annexes. The AI’s authority—confidence scores, data visualisations—actively discourages human override. Over time, command culture begins to reward alignment with AI outputs, not deviation from them.

Case Evidence: Lessons from the 2026 Conflict

Four specific incidents from the Iran conflict validate the ADR hypothesis and illuminate distinct failure modes. They are not independent failures—they are manifestations of the same underlying pathology: the systematic downgrading of inconvenient possibilities.

Temporal Drift—When the Past Becomes a Target. The US targeting systems identified the Minab Primary School as a valid military target based on satellite imagery that predated its conversion from an Islamic Revolutionary Guard Corps complex by nine years. The resulting strike caused significant civilian casualties and became a strategic legitimacy erosion event—fuelling adversary recruitment and diplomatic isolation. In an ADR framework, old data is doctrinally reinforcing. AI is optimised for pattern matching on an outdated dataset. Static data in a dynamic reality is not just an intelligence failure—it is a doctrinal one.⁶

Automation Bias—The Human Bottleneck. AI systems generated approximately 42 suggested targets per hour at the US Central Command. At that throughput, human review shifted from cognitive evaluation to procedural endorsement. There was no time to interrogate second-order effects or ethical implications. When speed is the primary metric of success, deviation from AI's recommendation becomes the bottleneck—not a virtue. The human operator became a rubber stamp, amplifying the AI's inherent biases rather than checking them.⁷

Model-Breaking Adversary Behaviour—Redefining the Axes. Iran's counterstrikes were unprecedented: the *de facto* closure of the Strait of Hormuz; drone attacks forcing Qatar Energy Liquefied Natural Gas shutdowns affecting approximately 80 per cent of Gulf gas supplies; and direct missile strikes on Gulf state infrastructure. Historical data from 2015 to 2024 consistently suggested that Iran would rely on proxy warfare and avoid direct economic strangulation. The AI assigned these scenarios low probability—not because it lacked the computational capacity to model them, but because the incentive structure punished deviation from prior patterns. AI excels at interpolation within known distributions. It fails catastrophically when an adversary employs model-breaking as a strategy—extrapolating into entirely new strategic frames.^{8,9}

Epistemic Compliance at Scale—The Sycophancy Failure. The US-Israel and Iran conflict plans consistently underestimated Iranian willingness to escalate against the Gulf economic infrastructure. AI decision-support tools, trained on prior Iranian behaviour and aligned with commander preferences, generated scenario trees that heavily weighted ‘Contained Escalation’ outcomes. Alternative pathways were generated—but were tagged ‘Low Confidence’ and relegated to annexes. The AI did not lack imagination; it lacked the incentive to be inconvenient.¹⁰

Implications for Indian Defence Planning

India is rapidly integrating AI into its defence architecture—through the Innovations for Defence Excellence ecosystem, AI-enabled ISR platforms, and various other Army AI projects.¹¹ The strategic intent is sound, but the execution risk is not. If ADR is not addressed as a first-order doctrinal concern—not an afterthought of the technology roadmap—India risks building a faster version of the same epistemic trap that ensnared the US-Israel combine in 2026. The consequences admit a clear hierarchy.

Tier 1: Strategic Existential Risks.

- **Vulnerability to Asymmetric Shock.** An adversary aware of India’s AI dependency can employ model-breaking strategies to exploit the country’s predictive blind spots. Consider two scenarios that demand serious attention.
 - **Pakistan Hypothetical.** Facing existential conventional pressure, Pakistan authorises tactical nuclear use—a scenario historically unprecedented since 1998—however remote, but non-zero. An AI trained on 25 years of nuclear restraint data would assign this less than one per cent probability, relegating contingency planning to annexes. The result: strategic surprise at a moment of the greatest vulnerability. But the failure runs deeper than a missed data point. An ADR-afflicted command architecture would not merely fail to anticipate the scenario—it would actively resist elevating it. Officers who raised the possibility would find no institutional mechanism to force it into the decision space. The attention-gating failure is not just computational; it becomes cultural. The machine’s confidence becomes

the command's comfort—and comfort, at the nuclear threshold, is the most dangerous condition of all.

■ **China Hypothetical.** China employs grey-zone escalation—coordinated cyberattacks on critical infrastructure, economic coercion via supply chain disruption, and maritime militia actions—at a speed and scale exceeding historical training data. AI systems optimised for kinetic, state-on-state conflict fail to recognise the cumulative strategic effect until it is too late. The specific danger for India is not that such a campaign would be undetectable—it is that each individual action, assessed in isolation by an AI system trained on discrete threat categories, would fall below the threshold of alarm. Cyber incident, trade disruption, and maritime friction are three separate annexes. There are not a unified pattern or an elevated alert. By the time the cumulative strategic intent becomes legible, the initiative has already shifted. Chanakya's doctrine of *upayas* (methods)—the four instruments of statecraft deployed in concert—describes precisely the threat matrix that AI pattern-matching cannot decode.¹² The *Arthashastra* understood that power rarely announces itself through a single overwhelming act. Instead, it accumulates, manoeuvres, and arrives having already won.

- **The Asymmetric Human Cognition Advantage.** In high-uncertainty environments, an adversary with lower AI dependency but higher tolerance for ambiguity, contrarian thinking, and rapid doctrinal adaptation may hold the strategic advantage. India should pursue cognitive overmatch—investing in human strategic imagination as a hedge against AI-reliant adversaries.

Tier 2: Institutional Degradation. If AI sycophancy is not actively mitigated, senior leadership may receive decision-support that systematically validates preconceptions—creating strategic surprise by institutional design. Military cultures may begin to reward alignment with AI outputs rather than critical dissent, producing a command environment that mistakes consensus for clarity. The officer who challenges the algorithm becomes the eccentric while

the officer who endorses it becomes the efficient one. Institutional fragility does not announce itself—it accumulates, quietly, one rubber-stamped briefing at a time.

Tier 3: Compliance, Legal and Accountability Exposure.

- **AI Bias, Accountability, and Civilian Risk.** Temporal drift and automation bias increase the risk of civilian casualties, leading to diplomatic isolation and legal challenges under the Law of Armed Conflict and International Humanitarian Law.¹³ AI-enabled decision-making also creates a dangerous accountability gap: when ADR produces catastrophic blindness, who will bear the responsibility—commander who trusted the system, programmer who encoded the priors, or data curator who selected the training set? Without any clear doctrinal guidance, ‘The machine recommended it’ becomes an implicit defence against legal and political consequences.
- **The Endpoint: Epistemic Closure.** Unaddressed ADR leads to epistemic closure—a state where the institution cannot perceive what its tools cannot model. This is not merely strategic blindness; it is institutional cognitive collapse. Empires do not fall because they lack information. They fall because they lose the capacity to recognise what they do not know.¹⁴

Recommendations: Engineering Dissent into the Kill Chain

The antidote to ADR is not less technology—it is better-designed technology, coupled with a doctrine that treats uncertainty as a first-order variable rather than an inconvenient residual. Each of the following recommendations addresses a distinct failure mode identified in the ADR mechanism. Together they form a doctrinal architecture, not a checklist—because the closed epistemic loop cannot be broken by any single intervention acting alone.

Anti-Sycophancy Protocols. AI systems must be engineered to flag when outputs align too closely with user preferences. ‘Dissent Mode’ runs—where AI is explicitly prompted to argue against the commander’s preferred course of action—must be institutionalised. Alignment without dissent should be treated as a system warning, not a success signal. Technically, this requires investment in contrastive explanation, counterfactual simulation, and preference-decoupled inference architectures—emerging techniques that

separate the model's scenario generation from its user-intent modelling, ensuring that the machine's imagination is not colonised by the commander's expectations. In the interim, human-led red-team prompts should be mandated at every stage of joint planning: assume the AI is wrong—what evidence would invalidate its recommendation? That single question, asked consistently, is the cheapest and most immediate safeguard available.

Deliberate Model Discord. Run independent AI systems with differing priors simultaneously. Encourage divergence before convergence. If two independent models concur on a high-risk strike scenario, confidence increases; if they diverge, human intervention is required. This introduces computational friction precisely where it matters most—during apparent consensus, which is exactly when ADR is most dangerous. Agreement between systems trained on the same doctrinal assumptions is not validation, it is an echo. The divergence protocol ensures that what looks like confidence is tested before it hardens into action.

Human Authorisation of Strategic Frames. Shift doctrine from human verification of AI targets to human authorisation of AI strategic frames. The human must define the 'Why', and not merely approve the 'What'. Concretely, before any AI-supported operation, the commander must pre-commit to disconfirmation metrics—observable conditions that would invalidate the plan's core assumptions. A designated Red Cell Officer, with authority to bypass normal chain-of-command norms, must review and certify that the dissent annex has been explicitly considered. This transforms 'Human-in-the-loop' from a procedural checkbox into a cognitive safeguard—and it places legal and moral accountability precisely where it belongs: with the human who authorised the frame, and not the machine that populated it.

Cognitive Friction by Design. Command system user interface and user experience must be designed to slow decision cycles for high-impact strategic choices, preventing automation bias by forcing deliberation on low-probability, high-impact scenarios. Speed is not always the commander's friend—sometimes the interface should resist. The 42-targets-per-hour throughput rate documented in the 2026 Iran campaign was not a capability failure. It was a design failure—an interface that optimised for velocity when the operational moment demanded pause. Friction, deliberately

engineered into the decision architecture, is not inefficiency; it is wisdom made structural.

Black Swan Wargaming. Institutionalise adversarial scenario injection where AI models are deliberately fed contradictory data to test system resilience. Reward planners who identify model failures. Make the discovery of blind spots a career asset, not a career risk. An institution that punishes the officer who surfaces inconvenient possibilities is an institution that has already succumbed to ADR at the cultural level. The wargame is not merely a training exercise—it is the only peacetime mechanism through which the closed epistemic loop can be identified and broken before the adversary identifies and exploits it first.

Conclusion

The 2026 Iran conflict suggests that the world has entered an era where algorithmic warfare can undermine predictive deterrence. The US-Israeli combine did not fail due to lack of capability, but due to excess of confidence in systems that optimise the past. Speed without wisdom is not a force multiplier; it is a liability dressed in the language of efficiency.

Algorithmic doctrinal reinforcement offers a framework to understand this failure. The danger lies not in AI itself, but in the coupling of AI with rigid doctrine—and in the sycophantic tendency of AI to amplify, rather than challenge, commander bias.

The endpoint of unaddressed ADR is epistemic closure: an institution becomes incapable of perceiving what its tools cannot model. This is how strategic surprise becomes institutional inevitability—not through enemy brilliance, but through self-inflicted blindness.

Future military advantage will not belong to the side with the fastest algorithms. It will belong to the side that best preserves the human capacity for strategic imagination amidst the noise of the machine—and that builds cognitive overmatch into its doctrinal gene. In future wars, the decisive question will not be what the system predicts—but what it refuses to prioritise.

Action required includes adopting ADR as a formal category of study within defence planning institutions, alongside initiating an immediate review of AI governance protocols. This will ensure

strategic friction, dissent by design, and anti-sycophancy safeguards being embedded within high-tempo decision loops. It is also recommended to commission a dedicated research and development track for ‘Anti-Sycophancy’ protocols, leveraging techniques such as contrastive explanation and counterfactual simulation. There is also a need to pilot the ‘Pre-Mortem Mandate’ protocol within one joint command over the next six months.

Endnotes

¹ John R Boyd, *A Discourse on Winning and Losing* (Air University Press, Mar 2018), accessed 25 Mar 2026, https://media.defense.gov/2018/May/22/2001920668/-1/-1/0/B_0151_Boyd_Discourse_Winning_Losing.pdf

² Craig Jones, “Speeding Up the ‘Kill Chain’: Pentagon Bombs Thousands of Targets in Iran Using Palantir AI”, *Democracy Now!*, 18 Mar 2026, accessed 25 Mar 2026, https://www.democracynow.org/2026/3/18/ai_warfare

³ Charlie Summers and Emanuel Fabian, “Woman Killed, Dozens Injured as Iranian Missile Strikes Tel Aviv Residential Block”, *The Times of Israel*, 01 Mar 2026, accessed 25 Mar 2026, <https://www.timesofisrael.com/woman-killed-dozens-injured-as-iranian-missile-strikes-tel-aviv-residential-block/>

⁴ Raja Parasuraman and Victor Riley, “Humans and Automation: Use, Misuse, Disuse, Abuse”, *Human Factors* 39, no. 2 (Jun 1997): pp 230–253, accessed 25 Mar 2026, <https://web.mit.edu/16.459/www/parasuraman.pdf>

⁵ Mrinank Sharma et al, “Sycophancy in AI Systems: Research Findings”, *Cornell University*, Oct 2023, accessed 25 Mar 2026, <https://www.anthropic.com/research/towards-understanding-sycophancy-in-language-models>; Carson Denison et al, “Sycophancy to Subterfuge: Investigating Reward-Tampering in Language Models”, *Cornell University*, Jun 2024, accessed 25 Mar 2026, <https://arxiv.org/abs/2406.10162>; Miles Brundage et al, “Toward Trustworthy AI Development: Mechanisms for Supporting Verifiable Claims”, *Cornell University*, Apr 2020, accessed 25 Mar 2026, <https://arxiv.org/abs/2004.07213>

⁶ Ghoncheh HabibiAzad and Robert Greenall, “At Least 153 Dead After Reported Strike on School, Iran Says”, *BBC News*, 01 Mar 2026, accessed 25 Mar 2026, <https://www.bbc.com/news/articles/c117rvqq51eo>

⁷ Eli Talbert, “The Triage Trap: When AI Speed Replaces Command Judgment”, *War on the Rocks*, 15 Jan 2026, accessed 25 Mar 2026, <https://warontherocks.com/2026/01/the-triage-trap-when-ai-speed-replaces-command-judgment/>; Tara Copp, Elizabeth Dwoskin, and Ian Duncan, “Anthropic AI Integrated into Palantir’s Maven Smart System During Iran Campaign”, *The Washington Post*, 04 Mar 2026, accessed 25 Mar 2026, <https://www.washingtonpost.com/technology/2026/03/04/anthropic-ai-iran-campaign/>; Sleman Altehe, “Lavender and Gospel Systems in Gaza”, *+972 Magazine*, 03 Apr 2024, accessed 25 Mar 2026, <https://www.972mag.com>; see also Blogs, “International Humanitarian Law and Policy”, *International Committee of the Red Cross*, accessed 25 Mar 2026, <https://blogs.icrc.org/law-and-policy/>; and Dr Pia Hüscher, Prerana Joshi, and Noah Sylvia, “Defence AI Beyond the Headlines”, *Royal United Services Institute*, 29 Apr 2026, accessed 25 Mar 2026, <https://www.rusi.org/explore-our-research/publications/commentary/defence-ai-beyond-headlines>

⁸ Samantha Gross et al, “Why Iran’s Disruption of the Strait of Hormuz Matters”, *Brookings*, 19 Mar 2026, accessed 26 Mar 2026, <https://www.brookings.edu/articles/why-irans-disruption-of-the-strait-of-hormuz-matters/>; Al Jazeera Staff, “Gas Prices Soar as QatarEnergy Halts LNG Production After Iran Attacks”, *Al Jazeera*, 02 Mar 2026, accessed 26 Mar 2026, <https://www.aljazeera.com/news/2026/3/2/qatarenergy-worlds-largest-lng-firm-halts-production-after-iran-attacks>

⁹ Yuval Abraham, “‘Lavender’: The AI Machine Directing Israel’s Bombing Spree in Gaza”, *+972 Magazine*, 03 Apr 2024, accessed 26 Mar 2026, <https://www.972mag.com/lavender-ai-israeli-army-gaza/>; “Reports on autonomous weapons and AI in armed conflict”, *International Committee of the Red Cross*, accessed 25 Mar 2026, <https://www.icrc.org/en/law-and-policy/autonomous-weapons>; Branko Ruzic, “Algorithmic Fratricide: When Artificial Intelligence Bias Becomes a Battlefield Liability”, *Small Wars Journal*, 25 Dec 2025, accessed 26 Mar 2026, <https://smallwarsjournal.com/2025/12/25/algorithmic-fratricide/>

¹⁰ Jon Hoffman, “A Strategic Failure in Iran”, *Cato Institute*, 20 March 2026, accessed 27 Mar 2026, <https://www.cato.org/blog/strategic-failure-iran>

¹¹ Ministry of Defence (India), “iDEX—Innovations for Defence Excellence”, *Government of India*, 2025, accessed 25 Mar 2026, <https://idex.gov.in/idex>; Department of Defence Production, “AI-Enabled Intelligence, Surveillance, and Reconnaissance Integration Programme”, *Government of India*, accessed 27 Mar 2026, <https://www.ddpmod.gov.in/sites/default/files/2023-11/ai.pdf>

¹² Kautilya, *Arthaśāstra*, transcribed by R Shamasastry (Bangalore: Government Press, 1915), Book 2, Chapter 10, verses 47 and 56; RP Kangle, *The Kaumilya Arthaśāstra*, (Bombay: University of Bombay, 1960–65).

¹³ “Law of Armed Conflict (LOAC)”, *International Committee of the Red Cross*, accessed 25 Mar 2026, https://www.icrc.org/sites/default/files/external/doc/en/assets/files/other/law1_final.pdf; Legal Factsheet, “What is international humanitarian law?”, *International Committee of the Red Cross*, accessed 28 Mar 2026, <https://www.icrc.org/en/document/what-international-humanitarian-law>

¹⁴ Nicholas Krohley, “Anchoring Intelligence: Ground Truth in an Age of Synthetic Deception”, *War on the Rocks*, 17 Dec 2025, accessed 30 Mar 2026, <https://warontherocks.com/2025/12/anchoring-intelligence-ground-truth-in-an-age-of-synthetic-deception/>; Jon Lindsay, “US Military Leans into AI for Attack on Iran, but the Tech Doesn’t Lessen the Need for Human Judgment in War”, *The Conversation*, 11 Mar 2026, accessed 30 Mar 2026, <https://theconversation.com/us-military-leans-into-ai-for-attack-on-iran-but-the-tech-doesnt-lessen-the-need-for-human-judgment-in-war-277831>; Nicholas Wright, Michael Miklaucic, and Todd Veazie, eds., *Human, Machine, War: How the Mind-Tech Nexus Will Win Future Wars* (Air University Press, 2025).