



The War Before the War: Predictive Intelligence and Human Control in Warfare

Lux Fulton

Introduction

In October of 1962, American reconnaissance aircraft returned from Cuba carrying photographs that eventually would become some of the most consequential images of the Cold War. The photographs revealed Soviet nuclear missile sites under construction only ninety miles from the coast of Florida. The United States objectively possessed enough intelligence to recognize the seriousness of the threat, yet President John F. Kennedy did not immediately authorize military action. Instead, these decisions were debated and diplomatic backchannels quietly opened. Nearly two weeks had elapsed before the crisis concluded through negotiated compromise.¹

The Cuban Missile Crisis illustrates an aspect of military decision-making that is frequently treated as a limitation rather than an asset. That feature is something as simple as time. Delays in intelligence analysis and command decisions can create obvious military disadvantages, but they can also provide opportunities to verify information and prevent an incorrect assessment from producing an irreversible military response. The relationship between speed and security is therefore not straightforward and has become increasingly debated as military technology has advanced.² Faster decision-making can provide an operational advantage while simultaneously reducing the opportunity for the human judgment intended to prevent escalation.

More than sixty years after the photographs taken in 1962, that tradeoff is becoming increasingly harder to preserve. Artificial intelligence (AI) is rapidly being integrated into environments such as military intelligence, surveillance, targeting, and decision-support systems precisely because it can process information and generate recommendations faster than human analysts can.³ Thus, AI is changing warfare not only by expanding what militaries can do, but by altering how consequential military decisions are reached. These systems are gradually placing algorithmic assessments closer to the center of decisions involving targeting and the use of force.⁴ Although military institutions, such as the U.S. Department of Defense and NATO, frequently emphasize that humans retain final decision-making authority, my paper argues that formal supervision does not necessarily demonstrate meaningful oversight.⁵

The Limits of Formal Human Control

The applications of AI in military technologies vary considerably in their characteristics and degree of autonomy. Some automate relatively narrow analytical tasks, while others use machine-learning models to identify patterns across large quantities of data and generate recommendations for human operators.⁶ Keeping a human "in the loop" therefore answers the

¹ Allison and Zelikow 1999.

² Morgan et al. 2020.

³ Dorsey and Bo, 2025.

⁴ Heller, 2023; Roff, 2014; Shany and Shereshevsky, 2025.

⁵ U.S. Department of Defense, 2023; NATO, 2021.

⁶ Cummings, 2017.

question of who formally authorizes an action, but not necessarily whether that person has the practical capacity to exercise independent judgment. If an algorithm processes information inaccessible to a human analyst, generates probabilistic conclusions through reasoning that cannot be fully reconstructed, or presents recommendations within an increasingly accelerated decision cycle, a human operator may possess the authority to say no without possessing sufficient opportunity to determine whether saying yes is justified. Human authorization and meaningful human control are therefore not necessarily equivalent.

AI-enabled decision-support systems have been rapidly transitioning however from theoretical military applications to tools actively being applied in contemporary conflicts. One example of this is during the Israel–Hamas war, where systems including Gospel and Lavender were used to produce targeting recommendations on an unprecedented scale. In the early weeks of the conflict alone, Lavender reportedly identified more than 37,000 potential targets.⁷ Although human operators formally retained responsibility for approving or rejecting targets before strikes occurred, reports still indicate that most recommendations were approved after as little as 20 seconds of review.⁸ AI-assisted targeting has also expanded into the Russia/Ukraine and United States/ Iran wars. Systems including Project Maven developed with companies such as Palantir, and incorporating technology from firms including Anthropic, have reportedly contributed to the generation of tens of thousands of potential strike targets.⁹ All together, these cases raise concerns that retaining humans formally within the decision process may not guarantee meaningful oversight, particularly when the speed and scale of algorithmic recommendations constrain independent review.

The increasing use of these systems has also intensified concerns about automation bias: the possibility that human operators may put excessive confidence in algorithmic recommendations and therefore exercise less independent judgment.¹⁰ Legal and policy discussions of military AI frequently identify automation bias as a threat to human control; however, there is emerging experimental research that complicates the assumption that military personnel will consistently defer to algorithmic recommendations. In 2026, Kahn, Horowitz, and Samotin, for example, found that West Point cadets demonstrated relatively calibrated trust in AI decision support rather than systematic automation bias.¹¹ Similarly, Shandler, Gross, and Shereshevsky found evidence of algorithmic aversion among Israeli military personnel, particularly in high-stakes scenarios.¹² What these findings suggest is that the problem is not simply that humans will trust AI too much, but that their reliance on algorithmic recommendations varies with context, experience, and also system design.¹³

There is broad agreement that meaningful human judgment should remain central to military decision-making, but far less agreement over what meaningful human control actually requires. Paul Scharre approaches this problem primarily through the design of human-machine relationships. Scharre argues that autonomy exists along a spectrum and that humans and

⁷ Anderson, 2025.

⁸ Layton, 2025.

⁹ Manson, 2026.

¹⁰ Dorsey and Moffett, 2025; ICRC 2024.

¹¹ Kahn, Horowitz, and Samotin, 2026.

¹² Shandler, Gross, and Shereshevsky, 2026.

¹³ Allen, Greg, and Taniel Chan. 2017.

machines can occupy different roles throughout the decision process.¹⁴ RAND research approaches the issue partly through decision speed, identifying AI's ability to reduce decision-making time as a potential military advantage while still acknowledging that accelerated timelines can also create strategic risks.¹⁵ Henry Kissinger, Eric Schmidt, and Daniel Huttenlocher focus more heavily on political institutions, arguing that governments will require new norms capable of adapting to the speed and complexity of AI.¹⁶ The approaches of these scholars differ greatly, but each argument nonetheless attempts to preserve human judgment through some combination of authority, system design, or institutional governance.

Each of these mechanisms is important, but those same mechanisms depend upon a more fundamental and underexamined condition. That condition is time. Authority matters only if a commander has sufficient opportunity to exercise it. Explainability matters only if an operator has time to evaluate the information a system provides. Likewise, political norms and diplomatic mechanisms can constrain escalation only if institutions have an opportunity to intervene before an algorithmic prediction has already generated a military response. This paper therefore defines deliberative time as the interval necessary for human decision-makers to independently evaluate an algorithmic recommendation, while also comparing alternative interpretations, verifying information when necessary, and intervening before an irreversible response occurs. This does not mean that military systems should simply operate more slowly, but when responding to an immediate physical threat, reducing latency may improve security. When interpreting probabilistic predictions about an adversary's future political intentions, however, speed can instead reduce the opportunity to distinguish evidence of hostile intent from uncertainty.

Comparison and Refutation of Alternative Approaches to Human Control

Scharre's argument provides the first test of this argument. Distributing authority between humans and machines can preserve formal human involvement, but it does not guarantee that humans have sufficient time to exercise that authority independently.¹⁷ This paper moves the debate beyond the simplistic question of whether a machine or a person "pulls the trigger," and recognizes that human involvement can occur at different stages of military action.¹⁸ Yet a commander may retain formal power to reject an algorithmic recommendation while lacking sufficient time to understand the consequences of acting upon it. Authority therefore asks whether a human can override a machine, while deliberative time would ask whether the human has a meaningful opportunity to decide whether they should. A system can satisfy the first condition while failing the second.

Another perspective can be analyzed through RAND. RAND cooperation identifies faster decision-making as one of AI's potential military advantages because processing information more quickly can allow a military to recognize and respond to a threat before an adversary can act.¹⁹ A version of this argument is that AI could actually increase meaningful human control

¹⁴ Scharre, 2018.

¹⁵ Morgan et al. 2020; Burdette et al. 2026.

¹⁶ Kissinger, Schmidt, and Huttenlocher, 2021.

¹⁷ Horowitz et al. 2018.

¹⁸ Scharre. 2018. 172-76.

¹⁹ Morgan et al. 2020; Burdette et al. 2026.

rather than diminish it. Earlier warning may give commanders more time to evaluate a threat, while also deliberately preserving additional decision time could sacrifice an operational advantage or allow a preventable attack to occur. In this argument, speed and deliberation therefore do not necessarily exist in opposition. There are clear circumstances in which this is clearly valuable. If sensors identify an incoming missile, shortening the time between detection and interception can increase the likelihood that the missile is destroyed before reaching its target.

This paper does not argue against that kind of speed, but it does argue that speed has a different effect when AI is being used to interpret uncertain political intentions. Hypothetically, a model predicting a specific probability that an adversary will initiate hostilities within a certain amount of time presents a different problem from a missile already in flight. The first requires humans to determine what incomplete evidence means and the second requires them to respond to an event already occurring. Reducing decision time in the second case may protect security. Reducing it in the first can remove opportunities to verify intelligence. The disagreement with a decision-speed approach is not whether faster decisions can be useful, but whether speed should be treated as an advantage independently of the kind of decision being accelerated.

Kissinger, Schmidt, and Huttenlocher, offer a third response where they argue that AI's speed and complexity will strain existing political institutions, requiring the development of new norms to preserve strategic stability.²⁰ This approach rejects the assumption that technological change automatically determines political outcomes. Governments can establish rules and institutional constraints controlling how emerging technologies are used. This paper agrees that such safeguards are necessary but argues that they operate only if political institutions have time to use them. Rules can determine what actions are permissible and political authorities can retain the power to interrupt military action, but none of these mechanisms, however, can exercise restraint retrospectively. The Cuban Missile Crisis is a historical but relevant example that illustrates how the existence of diplomatic channels mattered because President Kennedy and his advisers had time to use them.²¹ Intelligence could be debated before the United States committed itself to military action. A diplomatic channel existing on paper would have offered little protection if military responses had already occurred faster than political leaders could communicate through it.

This is more consequential when decision-support systems compress multiple stages of the process. Imagine that an AI system interprets an adversary's military movement as preparation for attack and recommends an immediate increase in readiness. The opposing state observes that increase, interprets it as threatening, and raises its own readiness in response.²² If those changes occur rapidly enough, diplomatic institutions may formally remain available while political leaders have progressively less opportunity to determine why the interaction began before responding to its newest stage. And, the failure is not only a lack of rules, authority, or communication channels, but also the element of time in which those safeguards could alter the sequence has been compressed. This compression is already visible at the

²⁰ Kissinger, et al. 2021

²¹ Allison and Zelikow 1999.

²² Burdette, et al. 2026.

tactical level in Ukraine, where AI-enabled and digital intelligence systems have been used to combine large quantities of battlefield information and that shorten the interval between identifying a potential target and acting upon that information.²³

In some cases, this process can occur within minutes, leaving substantially less time between intelligence interpretation and military response. The example from the Ukrainian war does not demonstrate that AI has already produced the interstate escalation sequence described above, but it demonstrates the mechanism on which that possibility depends.²⁴ Technological systems can compress decision cycles without formally removing humans from them. The concern fully arises when this same compression moves beyond solely identifying battlefield targets into interpreting an adversary's intentions. Time therefore occupies a different analytical position from authority, system design, explainability, norms, or even regulation. It is important to recognise that it does not replace these mechanisms, but it conditions whether they can operate. Authority figures determine who may decide. System design determines how humans and machines divide tasks. Explainability determines whether an operator can understand a system's recommendation. Norms and regulation determine what actors should or may do. Deliberative time determines whether humans have enough opportunity to use any of these safeguards before the decision becomes irreversible. Put more simply, the other approaches have asked how human control should be structured, but this paper asks what conditions must exist for that structure to function. Deliberative time is not another safeguard alongside authority, design, and regulation, but a condition that helps make those safeguards work.

The Rise of Predictive Warfare

Predictive AI threatens deliberative time not only because it processes intelligence faster, but also because it can move the point of military action earlier. Intelligence has always involved prediction, and military leaders rarely wait for complete certainty before acting. However, what changes with contemporary AI is the speed, scale, and apparent precision with which uncertain evidence can be converted into action.²⁵ A conventional intelligence assessment generally leaves decision-makers responding to observable developments, while predictive systems can encourage action based on probabilistic indicators of what an adversary may do next. Prediction therefore does not eliminate uncertainty, but instead changes when states are pressured to act on it.²⁶

The apparent precision of these statistical predictions intensify this problem. A model may identify patterns associated with previous military mobilizations without establishing that similar patterns represent the same political intentions in a new context.²⁷ An algorithm could be highly effective at identifying statistical correlations while also remaining incapable of determining whether an adversary is preparing an attack, conducting an exercise, responding defensively, or deliberately creating deceptive signals.²⁸ Once a probabilistic assessment enters

²³ Renault, 2026.

²⁴ Manson, 2026.

²⁵ Jervis, 1978.

²⁶ Allen & Chan, 2017.

²⁷ O'Neil, 2016.

²⁸ Pasquale, 2015.

military planning, commanders may face pressure to respond before that distinction can be resolved.²⁹

For example, consider a hypothetical scenario in the Taiwan Strait where an American AI-enabled intelligence platform monitors Chinese naval activity using satellite radar (SAR), maritime transponders (AIS), cyber telemetry, and electronic communications.³⁰ Rather than detecting ships departing from port, the system identifies small deviations from established baselines. These deviations might include changes in fuel procurement, maintenance schedules, or naval communications. None of these deviations independently indicates military action, but collectively the AI system calculates, for instance, an 88 percent probability that China intends to establish a partial maritime blockade within fourteen days. While hypothetical, this example mirrors the probability outputs produced by real-world AI systems.³¹ Military planners then face yet another problem because waiting for further confirmation risks leaving American forces out of position, but repositioning forces and raising readiness creates observable activity for opposing intelligence systems.

Suppose Chinese predictive systems then detect the increased American activity and interpret it as evidence of preparations for intervention. Chinese commanders increase their own readiness. American systems observe that response and update their original prediction upward. Neither system needs to malfunction and neither government needs to desire war. The issue is that actions taken in response to one prediction can create new evidence supporting the opposing side's prediction. The result is a potential feedback loop in which prediction participates in producing the behavior it was designed merely to anticipate.³²

Although public evidence does not yet establish a complete algorithmic feedback loop in Ukraine, Gaza, or other contemporary conflicts, these cases demonstrate the mechanism on which such a loop would depend; AI-enabled systems can substantially compress the interval between intelligence interpretation and military action. The Taiwan scenario extends this observable development into an interstate security dilemma, where one state's response to a prediction becomes new evidence for its adversary to interpret. Under these conditions, compression is itself consequential. As predictive decision cycles accelerate, a mistaken interpretation can produce observable changes in readiness or force posture before decision-makers have sufficient time to correct it. Those changes can then alter the adversary's assessment, transforming an initially uncertain prediction into behavior that appears to confirm it.

During the 1983 Soviet nuclear false alarm incident, the Soviet Oko early-warning satellite system mistakenly reported multiple incoming American Minuteman ICBM launches due to unusual high altitude sunlight reflections. Soviet military officers were placed in a position where automated, high confidence technological outputs created immense pressure to launch a retaliatory nuclear strike before political leaders could verify the reality. Fortunately, Lieutenant Colonel Stanislav Petrov recognized the system's anomaly and delayed reporting it as an

²⁹ Johnson, 2020.

³⁰ Allen and Chan, 2017 ; Horowitz et al. 2018.

³¹ CSET, 2024.

³² Jervis, 1978.

attack. This decision alone ultimately allowed human judgment to override machine inference.³³ The significance of the incident for contemporary AI is not that the Oko system resembled modern machine learning because it did not. Rather, the incident demonstrates the importance of preserving a human capacity to question apparently authoritative technological information before it produces an irreversible response, which is recognizable even in the year 1983.³⁴

Unlike traditional deterrence, where observable military actions typically precede political discussion, predictive warfare actually reverses the sequence. Instead, the statistical inference precedes physical action. Algorithmic expectations can generate the behaviors they predict, creating feedback loops in which opposing systems reinforce one another's assumptions before human decision-makers fully understand why military postures are changing.³⁵ A version of this dynamic has appeared in Gaza, where Lavender reportedly used patterns derived from individuals previously classified as militants to identify additional suspected targets, allowing earlier classifications to shape subsequent algorithmic assessments and military action.³⁶ That said, the issue is the treatment of uncertainty as objective knowledge. These outputs can create an illusion of certainty which therefore institutionalizes false confidence.³⁷ Thus, AI also accelerates belief by transforming uncertain predictions into actionable intelligence and altering the logic behind deterrence itself. The danger is not that AI predicts the future too quickly, but that military institutions may begin acting on an uncertain future before humans have had sufficient time to decide what that prediction actually means.

Who Controls the Predictive Infrastructure?

Deliberative time can support human control only if decision-makers have sufficient access and understanding of the systems they are expected to evaluate. Predictive military AI has increasingly depended on privately developed data and computational infrastructure. This complicates traditional assumptions about where strategic authority resides. Shoshana Zuboff's work demonstrates how digital systems derive power from collecting behavioral data to predict and shape behavior, while Kate Crawford emphasizes that AI systems reflect political and economic choices embedded in their design.³⁸ When applied to national security, this means that prediction is not simply a neutral process of identifying threats. System design influences which patterns become visible, which behaviors are classified as significant, and ultimately what information reaches human decision-makers.

The comparison between commercial and military prediction, however, only goes so far. That's because private platforms generally pursue commercial objectives, while military institutions operate under sovereign authority and in environments characterized by deception, uncertainty, and lethal consequences.³⁹ Yet governments still depend on private firms for elements of the computational infrastructure and AI capabilities supporting military operations.

³³ Hoffman, 2009.

³⁴ Scharre, 2018.

³⁵ Johnson, 2020.

³⁶ Anderson, 2025.

³⁷ Lee, 2018.

³⁸ Zuboff, 2019; Crawford, 2021.

³⁹ Benedikt, et al, 2020.

These firms do not acquire the legal authority to use force, but the systems they design can shape the information upon which sovereign decisions are made.⁴⁰

This is a further limit on deliberative time. Proprietary models and privately controlled infrastructure can restrict a government's ability to scrutinize how an algorithmic recommendation was produced. Additional time for review provides little protection if decision-makers lack sufficient access to the assumptions, data, or systems they are expected to evaluate.⁴¹ States may therefore retain formal authority over military force while depending on privately controlled infrastructure to determine what information enters the decision process.⁴² Therefore, human control requires not only time to deliberate, but sufficient access to make that deliberation possible.

The Black Box and the Limits of Human Judgment

If predictive warfare changes the manner in which states respond, the black box problem changes how well decision-makers can understand the predictions driving those responses.⁴³ The same complexity that gives these systems their predictive advantage can undermine human judgment. The more complex the model, the harder its conclusions may be to independently evaluate within the time available. Not all military AI is equally opaque. Rule-based systems and simpler models can retain relatively auditable logic. The more serious black-box problem is with deep neural networks whose outputs depend on interactions among vast numbers of learned parameters.

Frank Pasquale describes this phenomenon as the “black box”: systems in which users can observe inputs and outputs without meaningfully reconstructing the reasoning connecting them.⁴⁴ In military contexts, this opacity carries unusually high stakes. An incorrect commercial prediction may be audited or corrected after the fact, but an escalatory military decision may be irreversible. Explainable AI (XAI) attempts to address this problem by making model outputs more interpretable, but explanations of complex models can be approximations rather than complete accounts of how a prediction was produced.⁴⁵ More importantly, explanation requires interpretation, and interpretation requires time. In compressed military decision cycles, the time required to decode an AI recommendation is precisely the time that increased processing speed is placed under pressure.⁴⁶ If commanders lack either the information or time necessary to independently evaluate a recommendation, formal human authorization then provides little meaningful oversight.

In the case of an international crisis, governments ultimately act not on probabilities themselves, but on what they believe those probabilities mean. Determining that meaning correctly takes time to consider alternative explanations. As predictive systems compress that

⁴⁰ Bienvenue et al. 2025.

⁴¹ Pasquale, 2015.

⁴² Bienvenue, et al. 2025.

⁴³ Allen and Chan, 2017.

⁴⁴ Pasquale, 2015.

⁴⁵ Rudin, 2019.

⁴⁶ Pasquale, 2015.



time, they risk accelerating the demand for a decision without equally accelerating the human judgment necessary to make it.

Automation Bias and Institutional Deference

Research demonstrates that individuals can over-rely on automated recommendations, sometimes deferring to them despite contradictory information.⁴⁷ In national-security contexts, this reliance varies with factors including users' knowledge of AI, trust in the system, perceived competence, and task difficulty.⁴⁸ Military decision-making may intensify these concerns because commanders operate under pressure while AI systems are introduced precisely to process information beyond human capacity.⁴⁹ Additionally, when computational demands exceed what operators can independently evaluate, reliance on external systems can become a form of cognitive offloading. Challenging an apparently authoritative output requires precisely the attention and time that accelerated command systems seek to conserve.⁵⁰ Under these constraints, human control depends on whether operators retain the time, knowledge, and institutional capacity to challenge it.

The Collapse of “Diplomatic Friction”

Beyond individual decisions, these dynamics can intensify the classic Security Dilemma, which Herz defined as an environment where defensive actions by one state can unintentionally threaten another. In contrast, Jervis emphasized how uncertainty about motives makes military behavior especially vulnerable to misinterpretation.⁵¹ Historically, however, intelligence analysis and diplomatic communication created time between observing a potential threat and responding to it. This temporal buffer constitutes what this paper calls diplomatic friction: the delay that allows states to verify intelligence, reconsider assumptions, and communicate during crises. Although often treated as inefficiency, such friction can interrupt escalation before uncertainty becomes military action.⁵²

When opposing states rely on predictive systems, actions taken in response to uncertain assessments can rapidly become new signals for an adversary to interpret. The Security Dilemma can therefore become a feedback loop in which algorithmic assessments reinforce perceptions of aggression before political leaders can clarify intentions.⁵³ AI compresses the Security Dilemma not simply by accelerating decisions, but by reducing the interpretive space in which humans can question them.

Managing a Technology We Do Not Trust

The problem is not simply that military AI can fail, but that the qualities that make it strategically attractive can also make its failures more consequential. Unfortunately, identifying

⁴⁷ Parasuraman and Manzey, 2010.

⁴⁸ Horowitz and Kahn, 2024.

⁴⁹ Kahn, Probasco, and Kinoshita, 2024.

⁵⁰ Risko and Gilbert, 2016.

⁵¹ Herz 1950; Jervis 1978

⁵² Morgan, et al. 2020.

⁵³ Johnson, 2020.

this danger is easier than eliminating it. AI is already being integrated into military institutions, and states possess powerful strategic incentives not to abandon these capabilities unilaterally. The challenge therefore shifts from whether military AI should exist to how its most destabilizing characteristics can be governed. Managing predictive warfare requires accepting the political reality of continued AI development without treating unrestricted integration as either inevitable or desirable.

Addressing the Limits of Treaty-Based Regulation

One possible approach is an international agreement regulating military AI. More than 120 states have supported calls for a legally binding instrument regulating or prohibiting autonomous weapons operating without meaningful human control.⁵⁴ Yet several technologically advanced military powers, including the United States, Russia, and Israel, have opposed a comprehensive ban, while China has supported some form of legally binding prohibition.⁵⁵ Other governments maintain that existing international humanitarian law, including the Law of Armed Conflict (LOAC), can govern emerging autonomous weapons.⁵⁶ A comprehensive treaty therefore appears unlikely in the near term. States developing advanced military AI have substantial advantages in these capabilities and face incentives not to restrict themselves while competitors continue development.⁵⁷ Predictive warfare creates an additional regulatory problem in which systems capable of producing escalation do not necessarily qualify as autonomous weapons. An algorithm that predicts an attack can influence escalation without firing a weapon. A treaty focused exclusively on autonomous weapons could therefore regulate the final act of force while leaving the predictive infrastructure preceding it largely untouched.

If comprehensive prohibition remains politically difficult, states could instead pursue confidence-building measures (CBMs) that reduce risk without requiring them to abandon military AI. Transparency will remain limited by classified military research, but common standards for weapons review provide one starting point.⁵⁸ Article 36, for example, has advocated legal reviews that consider system predictability and reliability, opportunities for timely human intervention, and accountability.⁵⁹ Shared standards could therefore establish minimum expectations for preserving meaningful human judgment even among states unwilling to accept a ban.

International standards for testing and evaluation could serve a similar function. Discussions within the Convention on Certain Conventional Weapons have emphasized the need for rigorous testing so that military personnel can reasonably anticipate system behavior across intended operational environments.⁶⁰ Such standards could reduce the likelihood that misunderstood system behavior produces unintended military consequences. For predictive warfare, however, these safeguards must extend beyond weapons that autonomously select and engage targets. AI can contribute to escalation by predicting attacks or recommending

⁵⁴ Human Rights Watch, 2025.

⁵⁵ Buchholz, 2026 ; Human Rights Watch, 2025.

⁵⁶ United Nations, 2018.

⁵⁷ Allen and Chan, 2017.

⁵⁸ United Nations, 2018; Scharre, 2018.

⁵⁹ Roff and Moyes, 2016.

⁶⁰ United Nations, 2018; Morgan, et al. 2020.

force-posture changes while humans formally retain final authority.⁶¹ International standards should therefore evaluate not only who authorizes military action, but whether that person has sufficient time and information to independently assess the recommendation.

Political Commitment throughout the Development of AI

Where legal regulation remains unattainable, states still have the power to establish political commitments governing military AI. Unlike treaties, these declarations would not create binding legal obligations, but would establish a shared expectation of responsible behavior. The 2018 UN CCW Group of Governmental Experts, for example, affirmed that human responsibility for decisions involving weapons systems must be retained and that accountability cannot be transferred to machines.⁶² Although political commitments are typically easier to abandon than treaties, they are capable of establishing norms while creating a foundation for even stronger confidence-building measures. For predictive warfare, human control must require more than placing a human somewhere within the decision chain. The 1983 Soviet false alarm demonstrates the importance of preserving the capacity to question apparently authoritative technological information before it produces an irreversible response.⁶³ Political commitments should therefore protect not only human authorization but the deliberative delay necessary to make that authorization meaningful.

Conclusion

AI offers substantial military advantages while introducing new strategic risks. Preserving human control requires more than keeping a person formally within the decision-making process. If predictive systems operate faster than humans can evaluate their conclusions, then human authorization risks become a procedural formality rather than an independent decision. The central question is therefore not simply whether humans remain involved, but whether they retain sufficient time to exercise independent judgment.

International security should consequently treat deliberative time as a necessary condition of human control. This does not require slowing every military system, but rather requires distinguishing between situations in which speed enhances security and those in which it magnifies uncertainty. Automated interception of an observable threat may demand milliseconds; a prediction about another state's political intentions does not. In the latter case, human review is not inefficiency but a safeguard against escalation. The future of military AI therefore depends not only on how quickly machines can predict conflict, but on whether states preserve enough time for humans to prevent it. The war before the war is increasingly algorithmic, but whether prediction becomes conflict must remain a human decision.

⁶¹ Johnson, 2020; Horowitz and Kahn, 2024.

⁶² United Nations, 2018.

⁶³ Hoffman, 2009.



Bibliography

1. Abraham, Yuval. 2024. "‘Lavender’: The AI Machine Directing Israel’s Bombing Spree in Gaza." *+972 Magazine*, April 3, 2024.
2. Allen, Greg, and Taniel Chan. 2017. *Artificial Intelligence and National Security*. Cambridge, MA: Belfer Center for Science and International Affairs.
3. Allison, Graham, and Philip Zelikow. 1999. *Essence of Decision: Explaining the Cuban Missile Crisis*. 2nd ed. New York: Longman.
4. Bienvenue, Emily, Maryanne Kelton, Zac Rogers, Michael Sullivan, and Matthew Ford. 2025. "Private Tech Companies, the State, and the New Character of War." *Carnegie Endowment for International Peace*, December 1, 2025.
5. Buchholz, Katharina. 2026. "Twelve Countries Say No to Banning Autonomous Weapons." *Statista*, January 27, 2026.
6. Burdette, Zachary, et al. 2026. *How Artificial Intelligence Could Reshape Four Essential Competitions in Future Warfare*. Research Report. Santa Monica, CA: RAND Corporation.
7. Center for Security and Emerging Technology (CSET). 2024. *AI and National Security: Current Capabilities and Emerging Dilemmas*. Washington, DC: Georgetown University.
8. Crawford, Kate. 2021. *Atlas of AI: Power, Politics, and the Planetary Costs of Artificial Intelligence*. New Haven: Yale University Press.
9. Cummings, M. L. 2017. *Artificial Intelligence and the Future of Warfare*. Research Paper. London: Chatham House.
10. Herz, John H. 1950. "Idealist Internationalism and the Security Dilemma." *World Politics* 2 (2): 157–80.
11. Hoffman, David E. 2009. *The Dead Hand: The Untold Story of the Cold War Arms Race and Its Dangerous Legacy*. New York: Doubleday.
12. Horowitz, Michael C., Gregory C. Allen, Elsa B. Saravalle, Anthony Cho, Kara Frederick, and Paul Scharre. 2018. *Artificial Intelligence and International Security*. Washington, DC: Center for a New American Security.
13. Human Rights Watch. 2025. "UN: Start Talks on Treaty to Ban ‘Killer Robots.’" May 21, 2025.
14. Jervis, Robert. 1978. "Cooperation Under the Security Dilemma." *World Politics* 30 (2): 169–214.
15. Johnson, James S. 2017. "Automating the Military Decision-Making Cycle: AI and Great Power Competition." *Defense & Security Analysis* 33 (4): 312–27.
16. Johnson, James S. 2020. "Artificial Intelligence & Future Warfare: Implications for International Security." *European Journal of International Security* 5 (2): 147–69.
17. Kahn, Lauren, Emelia Probasco, and Ronnie Kinoshita. 2024. *AI Safety and Automation Bias: The Downside of Human-in-the-Loop*. Washington, DC: Center for Security and Emerging Technology.
18. Kissinger, Henry A., Eric Schmidt, and Daniel Huttenlocher. 2021. *The Age of AI: And Our Human Future*. New York: Little, Brown and Company.
19. Lee, Min Kyung. 2018. "Understanding Perception of Algorithmic Decisions: Fairness, Trust, and Emotion in Response to Algorithmic Management." *Big Data & Society* 5 (1): 1–16.

20. Morgan, Forrest E., Benjamin Boudreaux, Andrew J. Lohn, Mark Ashby, Christian Curriden, Kelly Klima, and Derek Grossman. 2020. *Military Applications of Artificial Intelligence: Ethical Concerns in an Uncertain World*. RR-3139-1. Santa Monica, CA: RAND Corporation.
21. NATO. 2021. *Summary of the NATO Artificial Intelligence Strategy*. October 22, 2021.
22. O'Neil, Cathy. 2016. *Weapons of Math Destruction: How Big Data Increases Inequality and Threatens Democracy*. New York: Crown.
23. Parasuraman, Raja, and Dietrich H. Manzey. 2010. "Complacency and Bias in Human Use of Automation: An Attentional Integration." *Human Factors* 52 (3): 381–410.
24. Pasquale, Frank. 2015. *The Black Box Society: The Secret Algorithms That Control Money and Information*. Cambridge, MA: Harvard University Press.
25. Renault, Clément. 2026. "AI, Intelligence Processing, Analysis, and Decision-Making in the Russo-Ukrainian War." *International Journal of Intelligence and CounterIntelligence*.
26. Risko, Evan F., and Sam J. Gilbert. 2016. "Cognitive Offloading." *Trends in Cognitive Sciences* 20 (9): 676–88.
27. Roff, Heather M., and Richard Moyes. 2016. *Meaningful Human Control, Artificial Intelligence and Autonomous Weapons*. London: Article 36.
28. Rudin, Cynthia. 2019. "Stop Explaining Black Box Machine Learning Models for High Stakes Decisions and Use Interpretable Models Instead." *Nature Machine Intelligence* 1 (5): 206–15.
29. Scharre, Paul. 2018. *Army of None: Autonomous Weapons and the Future of War*. New York: W. W. Norton.
30. United Nations. 2018. *Report of the 2018 Session of the Group of Governmental Experts on Emerging Technologies in the Area of Lethal Autonomous Weapons Systems*. CCW/GGE.1/2018/3. Geneva: United Nations.
31. U.S. Department of Defense. 2023. *DoD Directive 3000.09: Autonomy in Weapon Systems*. January 25, 2023.
32. Zuboff, Shoshana. 2019. *The Age of Surveillance Capitalism: The Fight for a Human Future at the New Frontier of Power*. New York: PublicAffairs.