

AI: Geopolitics and National Security

Sarah Kreps, Ben Buchanan, Michael Horowitz, and Erica Lonergan

Abstract: Artificial intelligence is reshaping geopolitics by transforming how states generate, interpret, and project power. As a general-purpose technology, AI scales perception and prediction across military and civilian systems while embedding geopolitical competition within global markets, infrastructure networks, and technological supply chains. This chapter argues that AI rivalry increasingly turns on control over material and organizational foundations – including semiconductors, compute infrastructure, and industrial coordination – rather than algorithmic innovation alone. At the same time, AI capabilities spread rapidly through commercial platforms, open-source ecosystems, and battlefield experimentation, creating persistent uncertainty about relative advantage and intensifying security dilemma pressures. Examining military integration, US–China competition, material foundations, cyber operations, and governance, the chapter shows how diffusion, infrastructural constraints, and institutional adoption shape technological competition. Because AI capabilities resist clear measurement and verification, governance is shifting toward indirect mechanisms – export controls, standards-setting, and multistakeholder coordination – rather than traditional arms-control regimes.

Disclosure: Generative AI tools were used during the drafting process for idea generation, organizational structuring, and minor language refinement. The authors retain full responsibility for the content, analysis, and interpretation presented in this work.

Introduction

Artificial intelligence is reshaping international security. Unlike earlier military technologies, it does not introduce a single new weapon or platform. Instead, it changes how states generate, interpret, and exercise power. Earlier technological revolutions commonly studied in international relations transformed specific dimensions of warfare. Gunpowder expanded the reach of force. Nuclear weapons multiplied destructive potential. Cyber capabilities extended intrusion into digital systems. Artificial intelligence operates differently. It functions as a general-purpose technology that scales perception, prediction, and decision-making—activities historically constrained by human cognition and organizational capacity. By accelerating information processing, expanding analytical reach, and enabling new forms of automation across military and civilian domains, AI alters how power is produced and applied rather than simply adding new instruments of force.

This chapter argues that although autonomous weapons have attracted the most attention in debates about military AI, the broader geopolitical consequences of AI arise from the infrastructure, industrial systems, and institutional processes required to deploy AI at scale. Advantage increasingly depends on the ability to mobilize compute infrastructure, coordinate complex technological ecosystems, and integrate AI into organizational decision-making. In this sense, competition over AI resembles competition over industrial capacity and technological infrastructure rather than a traditional arms race centered on discrete weapons systems.

A fundamental challenge with AI is that it is dual use. Many of the most advanced AI systems are developed by private firms operating in commercial technology sectors, while military applications often depend on integrating those capabilities into existing command structures and operational practices. As a result, AI rivalry increasingly unfolds through control over semiconductor supply chains, compute infrastructure, data centers, and standards governing technological ecosystems. These dynamics are particularly visible in the rivalry between the United States and China, where competition extends beyond research laboratories and battlefield applications toward the material and industrial systems that enable AI development.

These characteristics also complicate the ability of states to assess relative technological capability, an important component of evaluating military power and strategic advantage. Performance depends simultaneously on algorithms, data, computational infrastructure, and organizational adoption. Some of the most visible systems operate in commercial settings, while many military applications remain classified. At the same time, AI capabilities spread rapidly through open-source ecosystems, commercial spillovers, and battlefield experimentation. Concentration and diffusion therefore occur simultaneously, producing persistent uncertainty about relative technological advantage.

Uncertainty, in turn, intensifies strategic competition. When states cannot easily assess rivals' technological progress, they often assume worst-case scenarios. Technological developments that appear experimental or defensive may be interpreted as signals of offensive capability. Under these conditions, AI amplifies familiar security dilemma dynamics even when underlying shifts in material power remain ambiguous. Governments are therefore competing not only over military applications but also over the industrial infrastructure and technological ecosystems that support AI development.

Understanding these developments requires situating AI within broader theories of technological change and international competition. Previous research shows that certain technologies transform global power distributions by imposing organizational and financial demands that only some states can meet (Posen 1984; Horowitz 2010). AI shares these characteristics but introduces an additional complication: technological advantage depends not only on invention but also on institutional adoption, infrastructure coordination, and diffusion pathways that spread capabilities across a widening range of actors. The effects of AI therefore unfold unevenly across domains and states, producing layered patterns of concentration and diffusion that complicate traditional measures of military power.

This chapter examines how AI affects international security across four interrelated arenas. First, it explores how AI alters battlefield decision-making, military organization, and escalation dynamics through decision-support tools, autonomous systems, and doctrinal adaptation. Second, it analyzes how AI transforms great-power competition by shifting rivalry toward the industrial infrastructure and technological supply chains that sustain frontier development, with a focus on the US–China relationship. Third, it examines the material foundations of AI capability, focusing on semiconductors, compute infrastructure, and the political economy of scaling advanced technological systems. Fourth, it analyzes how AI interacts with cyberspace and information warfare by transforming intelligence analysis, expanding opportunities for deception, and reshaping the offense-defense balance in digital conflict.

These dynamics also create significant governance challenges. Traditional arms-control frameworks rely on clear weapons categories, centralized production, and observable verification mechanisms. AI lacks these characteristics. Governance therefore emerges through indirect and fragmented approaches—including export controls, technical standards, voluntary industry coordination, and alliance-based cooperation—that attempt to shape technological trajectories rather than prohibit specific capabilities.

Finally, AI diffusion expands the range of actors capable of influencing international security. Commercial models, cloud infrastructure, and open-source software lower barriers to entry for military and intelligence innovation in many applications outside the frontier of model development. As a result, states increasingly compete not only with rival governments but also

with technologically empowered private firms, insurgent organizations, and transnational cyber groups.

AI on the Battlefield and in Geopolitics

AI and the Reconfiguration of Military Power

Technological change has long shaped the distribution of power in international politics. International relations scholarship has shown that certain innovations transform the underlying requirements for generating and projecting military force. Horowitz (2010) describes nuclear weapons and aircraft carriers as “system-altering” capabilities because they imposed steep organizational and financial barriers that only a limited number of states could overcome. Similarly, Buzan and Herring (1998) identify “transformational” military technologies as those that fundamentally restructure security competition by altering the cost, scale, or character of warfare.

Artificial intelligence increasingly appears to belong within this category, but for reasons that differ in important respects from earlier military revolutions. Unlike nuclear weapons, missiles, or conventional armaments, AI does not derive strategic significance primarily from destructive force. Instead, AI functions as a dual-use, general-purpose technology that operates across civilian and military sectors simultaneously. Advances in data analytics, machine learning, and automated decision-support systems improve intelligence processing, logistics optimization, targeting, cyber operations, and command and control rather than directly increasing explosive capability. In this sense, AI reshapes military power less by replacing traditional weapons and more by transforming how states organize, integrate, and employ them.

This dual-use character complicates both strategic competition and governance. Because AI development is deeply embedded in commercial technology ecosystems, advances in civilian computing infrastructure, software engineering, and data management frequently translate into military capability. Scholars increasingly emphasize that AI’s integration across civilian and defense sectors has the potential to alter patterns of geopolitical competition by generating new dependencies, creating novel vulnerabilities, and shifting the locus of military advantage toward states that can mobilize technological ecosystems rather than simply field superior weapons systems (Horowitz 2018; Vaynman and Volpe 2025). As a result, AI’s strategic significance emerges not only in battlefield applications but also in global supply chains, export control regimes, alliance networks, and industrial policy debates.

The military and geopolitical effects of AI therefore operate through systemic and cumulative mechanisms rather than singular technological breakthroughs. AI may not transform warfare through a single decisive weapon comparable to nuclear arms, but its integration across sensing, analysis, and operational coordination has the potential to reshape how states generate and

sustain military power over time. Understanding AI's strategic impact requires attention not only to autonomous systems or weapons platforms, but to the broader technological, economic, and organizational ecosystems through which military capability is produced and exercised.

Horizontal Diffusion and Substitution Effects

In the context of national defense, AI use cases principally fall into three buckets: back-office administration, surveillance and decision support, and the battlefield. While back-office applications shape long-run military efficiency, the dynamics most salient for signaling, escalation, and competition emerge in decision support and battlefield integration, with the latter generating the greatest disruption to established military decision-making and raising unresolved questions about the scale and consequences of delegating lethal force.

Indeed, the back office uses of AI for militaries look much as they do in other parts of government or the private sector – optimizing personnel processes, acquisition workflows, peacetime logistics, and force management. While these applications appear mundane, their strategic significance lies in their cumulative effects: Improved administrative efficiency reduces friction in force generation, accelerates mobilization, and allows militaries to sustain complex operations with smaller bureaucratic and labor footprints. Although such uses carry risks if systems are poorly trained or deployed, their symmetry with civilian applications makes them less distinctive from an international relations perspective (Global Commission on Responsible Artificial Intelligence in the Military Domain 2025).

A second area where militaries are integrating AI is decision support, computational tools that use machine learning and large-scale data analysis to generate predictions, recommendations, or risk assessments that inform rather than replace human decision-making in complex operational environments (US Department of War 2026). These systems typically aggregate inputs from multiple sensors and data streams, apply machine-learning models to filter, rank, and prioritize information, and present assessments or recommended courses of action to human operators within existing command-and-control structures. By reducing cognitive load and accelerating information processing, decision support tools reshape how commanders perceive and act on the battlefield.

The most prominent domain for such systems is intelligence, surveillance, and reconnaissance, where AI can process vast quantities of data more quickly than human analysts and identify patterns that would otherwise remain obscured. Decision-support tools also increasingly assist commanders in operational planning and, in some cases, targeting decisions (Nadibaidze, Bode, and Zhang 2024). These applications reconfigure military power by shifting advantage toward actors able to integrate data, analytics, and command structures effectively. Over time, this transformation may redistribute strategic advantage away from states that rely primarily on

platform superiority or force size and toward those capable of sustaining complex data and algorithmic ecosystems.

These AI systems also compress decision timelines by accelerating information processing and operational planning cycles, potentially narrowing opportunities for deliberation during crises. Recent U.S. military operations in Iran illustrate how this dynamic has appeared in practice, with Anthropic's Claude AI integrated into tools like the Maven Smart System to synthesize intelligence reporting, propose hundreds of potential targets, prioritize them by importance, and provide precise location coordinates, all within highly compressed planning timelines that enabled over 1,000 strikes in the first 24 hours (Copp, Dwoskin, and Duncan 2026).

The systems raise concerns about automation bias – situations in which operators defer to algorithmic outputs beyond their demonstrated reliability. Experimental evidence suggests such bias may influence national security decision-making, particularly when users lack sufficient familiarity with the systems they employ (Horowitz and Kahn 2024). As reliance on AI-mediated analysis expands, such biases may introduce new pathways for misperception and escalation, especially in fast-moving conflict environments where verification and human oversight become more difficult to sustain.

The final category of military AI use lies on the battlefield itself in the context of lethal strike decisions. These dynamics are most visibly unfolding through drone warfare – understood here as the use of remotely piloted, semi-autonomous, and increasingly autonomous unmanned aerial systems for surveillance, targeting, and strike operations – which provides the clearest empirical illustration of how autonomy is being integrated into military operations.

Most deployed systems to date – including loitering munitions such as the Israeli HAROP, which function as hybrids between drones and missiles that loiter while searching for targets before striking – delegate navigation or target recognition while preserving human decisions over lethal force.

Fully autonomous strike systems remain difficult to field when militaries seek reliable target identification under adversarial conditions while minimizing collateral damage and complying with international humanitarian law. Achieving this standard requires resilient operation in contested communications environments, real-time onboard computation within strict size and power constraints, and algorithmic decision-making capable of distinguishing lawful targets from civilians under conditions of uncertainty. However, these systems become considerably easier to deploy if operators are willing to tolerate higher rates of false positives or accept reduced human oversight. The resulting trade-off highlights why autonomous strike capabilities remain both technologically and politically contested. They raise persistent questions about whether algorithms can reliably perform the subjective and high-stakes judgments required for lawful and proportional targeting decisions (Kaag and Kreps 2014).

These concerns are closely tied to ongoing debates about what constitutes meaningful human control over autonomous weapons systems. Department of Defense guidance emphasizes that autonomous weapon systems must allow commanders and operators to exercise appropriate levels of human judgment over the use of force, reflecting broader legal expectations that attacks remain subject to precaution and discrimination requirements even when mediated by advanced technologies (U.S. Department of Defense 2023; ICRC 2014).

Nonetheless, if faced with a situation where an adversary deploys an autonomous system and perceives advantages of doing so, a state may feel pressure to reciprocate. Those advantages may be that they reduce the time of decision loops – from identifying to striking time-sensitive targets – or circumvent the need for real-time communication if electronic warfare has been disruptive. The increasing reliance, however, creates new risks. Adversaries can deploy decoys, spoof sensors, or deliberately mix civilian and military assets to manipulate automated targeting systems. Rapid machine-driven strikes can also accelerate escalation if autonomous systems misidentify targets or launch follow-on attacks without human reassessment. These dynamics illustrate how time pressure, survivability concerns, and competitive adaptation can push militaries toward offensive lethal autonomy despite persistent uncertainty about legal compliance, escalation control, and accountability.

Diffusion, Assessment, and the Security Dilemma

Because military AI adoption unfolds in a competitive environment, states interpret adversary investments in autonomy not only as technological developments but also as potential shifts in relative military advantage. Even defensive or efficiency-driven deployments can generate uncertainty about adversaries' intentions and capabilities, particularly when autonomous systems reduce decision timelines or operate in communication-denied environments. Under such conditions, efforts to preserve deterrence credibility or avoid operational disadvantage can drive reciprocal adoption of AI-enabled systems, even when states recognize the risks of escalation or accidental conflict. These interactive dynamics shift the problem from one of technological capability alone to one of strategic competition under uncertainty.

The discussion raises the question of what it means to fall behind, which would be at the heart of understanding the consequences of AI for geopolitical competition and the balance of power. Unlike nuclear arsenals or missile inventories, AI resists simple quantification because there is more than one AI. Horowitz argues (2018) that AI is better understood as a general-purpose technology – akin to electricity or the combustion engine – whose significance depends on how organizations adopt and apply it. Should capability be judged by algorithmic novelty, model scale, compute volume, or integration into command-and-control systems? The absence of clear thresholds complicates both competitive assessments and governance frameworks. Visible private-sector deployments like ChatGPT capture public attention, while public-sector classified systems remain obscured, widening the gap between perception and reality.

A world characterized by limited information and uncertainty about the trajectory of AI development – particularly whether advances could lead to artificial general intelligence or superintelligence – creates incentives for worst-case thinking. If AI development includes a point at which recursive self-improvement produces rapid capability expansion, achieving that threshold first could plausibly confer significant strategic advantages. Even if such a discontinuity never materializes, the possibility that it might exist can influence state behavior. Importantly, the strategic consequences of AGI need not depend on precise technical definitions or even on its eventual realization; perceptions that such a capability threshold might exist can themselves drive competitive behavior. Because private firms, rather than governments, operate at the cutting edge of AI development, governments may struggle to forecast technological progress more than they have with traditional military systems. Under these conditions, states may interpret ambiguous AI advances by rivals as signals of potential offensive advantage, increasing the risk of security dilemmas.

As Robert Jervis (1978) argued in his seminal work on the security dilemma, states often respond less to what their rivals can do than to what they believe those rivals intend to do, and misperceptions in this domain can generate spirals of mistrust even in the absence of aggressive aims. Historical precedents illustrate how emerging technologies can intensify these dynamics. For example, the Reagan administration's Strategic Defense Initiative, though framed as a defensive effort to protect against nuclear attack, was widely interpreted by Soviet leaders as a potential attempt to neutralize deterrence, contributing to intensified strategic competition despite significant uncertainty about the system's technical feasibility.

More recent work by Henry Farrell and Abraham Newman (2019) extends this insight to an era of economic and technological interdependence, showing how control over networks, standards, and chokepoints can be interpreted simultaneously as defensive safeguards and as instruments of coercion. In this context, policies that are framed domestically as prudent risk management may be read abroad as hostile signals, reshaping expectations and accelerating rivalry even when underlying material balances change little.

The risk of hacking and spoofing algorithms, and the potential for algorithms to make decisions faster than humans can effectively review them, could further increase escalation risks (Johnson 2021). Especially for states that fear the disruption or destruction of command-and-control leadership structures – even in non-nuclear contexts – the prospect of warfare at machine speed may encourage reciprocal adoption of AI-enabled decision making. Operational control by AI systems, particularly when algorithms are optimizing weapons selection and targeting, could generate escalation dynamics that become difficult to manage or reverse (Johnson 2022).

Alternatively, some critiques of the spiral model specifically suggest that concerns about AI-enabled escalation are overstated. Reiter (2025) argues that the desire of leaders for control over escalation pathways means that there are more offramps for war than the spiral model

would suggest. Essentially, fear of spiral model-type scenarios leads leaders to tightly hold the reins of military escalation, which, in turn, decreases the risk that a conflict spirals out of control. He suggests that the desire of leaders for control of the escalation ladder means that leaders will be unlikely to delegate control of important military operations or decisions to artificial intelligence (Reiter 2025).

The US-China AI Competition

Nowhere is the interaction among technological uncertainty, competitive signaling, and concern about escalation control more salient than in the context of the U.S.–China rivalry. It reflects both the uncertainties surrounding technological adoption and employment and the broader structural competition accompanying China’s economic and military rise.

This section examines how competition over AI is shaped by material constraints, network position, and domestic political economy. It first contrasts the distinct visions through which the United States and China pursue AI leadership. It then shows that success at the frontier depends on access to a narrow set of material inputs – most notably compute and the infrastructure required to deploy it at scale. It also explains why these inputs function as strategic chokepoints that generate leverage but not full control within global production networks. The section next examines how domestic political economy conditions each country’s ability to mobilize and govern this infrastructure. Finally, it links these dynamics to the scaling logic of modern AI systems and to the semiconductor supply chains that anchor both the possibilities and limits of state influence.

Visions of AI Global Leadership

The United States and China signal the stakes of contemporary AI competition by treating leadership at the technological frontier as a strategic priority. In early 2026, China President Xi Jinping met with his advisors and called AI “epoch-making technological transformation” that required a “whole-of-nation” approach to develop domestic capabilities to assert international leadership (Udinmwun 2026). The United States’ core AI document is called “Winning the AI Race: America’s AI Action Plan” that suggests “Winning the AI race will usher in a new golden age of human flourishing, economic competitiveness, and national security for the American people” (White House 2025). The National Security Commission on AI (2021) writes that, “We know that whoever translates AI developments into applications first will have the advantage. Now we must act.”

Framing AI development as a race toward dominance risks reinforcing competitive pressures. Security dilemma dynamics help explain why the rhetoric of an AI arms race persists despite the fragmented and multidimensional nature of AI development. When states portray AI leadership as a zero-sum contest, they encourage adversaries to interpret technological progress as evidence

of offensive intent rather than defensive modernization. This perception can generate reciprocal investments in AI capability, even when advantage in one domain does not translate into overall technological dominance. In this sense, narratives of AI supremacy can operate as self-reinforcing signals that intensify rivalry, accelerate investment, and shift competition upstream toward control over the material and industrial foundations of AI development.

This interpretation reflects broader arguments that AI competition is unlikely to produce decisive technological dominance by any single state. Instead, leadership is fragmenting across multiple domains, including frontier model development, control over computing infrastructure, influence over global technological standards, and the integration of AI into physical and military systems (Kahl 2026).

If AI competition unfolds across multiple domains rather than along a single technological dimension, the question becomes what conditions the prospects for success across those domains. Fragmented leadership does not eliminate structural constraints. It shifts attention to the material systems that enable innovation at scale. Understanding the trajectory of U.S.–China AI competition therefore requires moving beyond rhetoric and institutional design to examine the infrastructural foundations that make frontier development possible.

Material Foundations of AI

Achieving leadership in AI depends on more than strategic intent or institutional design. It increasingly reflects control over the material and infrastructural foundations that underpin technological development. Insights gleaned from international political economy help explain why these foundations matter strategically. As Farrell and Newman (2019) argue, power in globalized technological systems often arises from dominance over key nodes within production and governance networks rather than direct control over outcomes. Actors positioned at these hubs can shape the technological options available to others by structuring access to essential inputs, production processes, and infrastructure. Applied to artificial intelligence, this framework directs attention to the components of AI development that exhibit high centrality, high switching costs, and concentrated governance. These characteristics increasingly define large-scale compute capacity, semiconductor manufacturing, and the infrastructure required to train and deploy advanced AI systems.

Among these inputs, compute has emerged as the binding constraint on AI development. Training and operating state-of-the-art models requires vast quantities of specialized hardware, dense networking capacity, reliable energy supply, and sustained capital investment. Unlike software or algorithmic innovations, which can diffuse relatively quickly, the physical infrastructure that underpins AI development depends on capital-intensive and slow-turnover production systems. Advanced chips must be designed, fabricated, and packaged through highly specialized manufacturing processes. Data centers must secure land, electricity generation,

cooling systems, and network connectivity. Large-scale model training requires simultaneous access to all of these inputs over extended time horizons. As a result, improvements in AI capability increasingly track the ability to mobilize and coordinate physical infrastructure rather than isolated breakthroughs in model architecture.

The semiconductor industry sits at the core of this material foundation. Control over advanced semiconductors, fabrication capacity, and the specialized equipment required to manufacture them has become a central constraint on technological leadership (Miller 2022). Leading-edge chips are produced through manufacturing processes that only a small number of firms can execute at scale, using lithography and fabrication tools that themselves rely on highly complex and geographically concentrated supply chains. These characteristics make advanced semiconductors difficult to substitute and slow to reproduce, even with substantial financial investment.

Compute therefore functions not merely as an input to AI development but as a structural foundation that shapes the pace and direction of innovation. States and firms that command sustained access to advanced chips and the infrastructure required to deploy them gain the ability to train larger models, iterate more rapidly, and push performance boundaries outward. Those that lack such access face constraints on scale, training frequency, and deployment capacity, although advances in algorithmic efficiency and model optimization can partially mitigate these limitations.

Advances by Chinese AI developers, most notably the release of the DeepSeek models, illustrate both the potential and limits of cost-efficient innovation at the AI frontier. DeepSeek reportedly achieved competitive performance on several benchmark tasks using significantly lower training costs and fewer computational resources than comparable Western frontier models, demonstrating that engineering optimization, data efficiency, and algorithmic refinement can partially offset raw compute scale. These developments suggest that high-end computing infrastructure is not the sole pathway to advancing model performance and that resource-constrained actors can narrow capability gaps through software and training innovations.

At the same time, evidence from model evaluation and deployment patterns indicates that cost-efficient models often remain less reliable on complex reasoning tasks, less scalable across diverse enterprise and multimodal applications, and slower to iterate at frontier performance levels. High-compute models retain advantages in training speed, robustness, and the ability to support large-scale commercial deployment. Rather than eliminating the importance of compute, the DeepSeek experience suggests that AI competition increasingly involves a dynamic interaction between efficiency gains and infrastructure scale, in which optimization lowers entry barriers but sustained access to large compute clusters continues to define the outer limits of performance and iteration speed.

These material dependencies generate geopolitical significance not just because they are scarce, but because they occupy central positions within global production and governance networks. Control over compute infrastructure, semiconductor manufacturing, and data-center deployment functions as a chokepoint in the development of advanced AI capabilities. These inputs sit upstream of model development and deployment, function as bottleneck resources across national AI ecosystems, and remain governed by a small number of firms and jurisdictions. States that control access to these nodes can raise costs, slow diffusion, and redirect competitors' development pathways without directly suppressing innovation or relying on overt coercion. Conversely, states that depend on infrastructure governed by external actors face structural vulnerability even when they retain technical expertise and domestic research capacity.

Chokepoint power in AI operates gradually and indirectly. Restrictions on access to advanced chips or fabrication tools rarely halt technological progress outright. Instead, they reshape the conditions under which progress occurs by influencing the scale, speed, and direction of innovation over time. Competitors may pursue alternative hardware architectures, increase investment in efficiency-oriented research, or redirect development toward applications that require less computational intensity. In this sense, structural power in AI resembles other forms of economic statecraft in which influence arises through the structuring of technological opportunity rather than the prohibition of technological development itself.

Understanding AI competition therefore requires attention not only to who develops new models, but also to who governs the infrastructure and production networks through which advanced AI becomes possible. Leadership in artificial intelligence increasingly reflects the capacity to control, mobilize, and scale material systems that shape technological trajectories across both military and commercial domains.

Scaling, Compute, and the Industrial Foundations of AI Power

The geopolitical significance of compute and semiconductor supply chains reflects a deeper technological dynamic within contemporary AI development. Progress at the frontier increasingly follows a scaling logic in which performance improvements derive from training larger models on greater quantities of data using expanded computational resources rather than from discrete conceptual breakthroughs. Under this logic, AI system performance improves as developers increase model size, training data, and compute allocation, even when underlying architectural innovations remain relatively modest. Countries such as Saudi Arabia and the United Arab Emirates are investing heavily in data centers and compute, using oil wealth to create a powerful AI hub in the Middle East (Satariano and Moser 2025).

This pattern appeared in earlier machine-learning research but became especially visible with the emergence of large language models. OpenAI's GPT-3, released in 2020, relied on design principles similar to earlier transformer architectures but expanded model parameters and

training data by orders of magnitude. The resulting gains in language generation and reasoning ability demonstrated that expanding computational scale had become a primary driver of AI capability. Subsequent model development across the industry has largely reinforced this scaling trajectory.

Scaling dynamics generate strong incentives toward industrial concentration. Training advanced models requires access to specialized chips, massive data-center capacity, and sustained electricity supply, all of which exhibit high fixed costs and strong economies of scale. These requirements sharply limit the number of actors capable of competing at the frontier and reinforce the importance of firms' and states' abilities to mobilize large-scale capital investment and infrastructure deployment. As a result, AI development increasingly resembles other capital-intensive industrial sectors in which performance gains depend on the ability to expand and coordinate physical production systems rather than solely on intellectual or algorithmic innovation. Whether AI leadership will prove strategically decisive is unclear, but the belief that it might do so could drive states toward intensified technological rivalry.

The semiconductor ecosystem illustrates how scaling logic translates into industrial concentration. Leading AI chips depend on advanced fabrication processes that only a small number of firms can execute at scale. Taiwan Semiconductor Manufacturing Company (TSMC) occupies a uniquely central position in global semiconductor production as the dominant contract manufacturer of advanced logic chips, producing processors designed by leading AI and computing firms and enabling the fabless production model that underpins much of the modern semiconductor industry.

Critical upstream suppliers such as ASML, the sole commercial supplier of extreme ultraviolet (EUV) lithography systems, provide the specialized equipment necessary for cutting-edge manufacturing. EUV lithography enables the production of advanced transistors but requires extraordinary engineering precision and integrates supply chains spanning multiple countries, reinforcing barriers to entry in semiconductor production. The long development timelines and capital requirements associated with EUV systems create substantial barriers to entry, reinforcing concentration in semiconductor production and limiting the speed with which new competitors can emerge.

China's effort to build an indigenous semiconductor industry highlights the depth of the constraints embedded in advanced chip manufacturing. Beginning in the mid-2010s, Chinese leaders identified dependence on foreign semiconductors as a strategic vulnerability and launched a state-directed push to indigenize production through massive public investment, most notably via the "Big Fund." The "Big Fund," formally the China Integrated Circuit Industry Investment Fund, is China's primary state-backed financing mechanism designed to build a domestic semiconductor industry through large-scale public investment. As Miller (2022) documents, Chinese semiconductor initiatives have faced persistent difficulty replicating the

precision manufacturing, supply-chain coordination, and tacit engineering expertise required for leading-edge chip production despite sustained funding and political backing. They remained dependent on foreign fabrication tools and equipment and, even when access remained available, they failed to achieve yields comparable to those of established producers. The experience highlights the long time horizons and organizational complexity involved in semiconductor production, suggesting that manufacturing capability cannot be rapidly generated through capital investment alone.

Similar challenges appear even among advanced economies. Despite extensive incentives and close partnerships with TSMC, the United States has struggled to replicate the full manufacturing ecosystem and yield performance of TSMC's operations in Taiwan, including at facilities located on U.S. soil (Goodman 2025). This pattern reinforces the extent to which advanced semiconductor production depends on deeply embedded industrial coordination, workforce specialization, and accumulated tacit knowledge. The reported use of commercial frontier models such as Anthropic's Claude within U.S. military intelligence systems during operations related to Iran illustrates how rapidly such commercial capabilities can migrate into operational military workflows (Copp et al. 2026).

Scaling dynamics do not produce absolute technological lock-in. Efficiency-oriented innovations can partially offset infrastructure constraints. Chinese AI firms, for example, have demonstrated the ability to produce competitive large language models through cost-efficient training and algorithmic optimization strategies. Such developments suggest that constraints on compute often redirect innovation toward alternative training approaches, cost reduction strategies, and model efficiency improvements rather than halting capability development altogether. However, scaling advantages remain significant because large compute clusters enable faster training iteration, broader multimodal integration, and more reliable deployment across complex commercial and military applications.

The scaling logic of AI therefore reinforces the strategic importance of compute infrastructure while simultaneously generating ongoing competition between scale and efficiency. This dynamic helps explain why AI development increasingly concentrates within a small number of firms and states while remaining open to partial diffusion through optimization and engineering innovation.

Semiconductors as Statecraft

The concentration of advanced semiconductor manufacturing and compute infrastructure has created new opportunities for states to treat technological supply chains as instruments of geopolitical strategy. Rather than functioning solely as commercial inputs, semiconductors increasingly operate as tools of economic statecraft through which governments attempt to shape the trajectory of technological development in rival states. In the context of artificial intelligence,

where access to large-scale compute capacity directly influences model training, deployment, and iteration speed, control over semiconductor supply chains has become a central lever of strategic competition.

By the late 2010s, the United States had begun to recognize advanced chip manufacturing as both a source of competitive advantage and a potential instrument of coercive leverage. China's repeated difficulties in indigenizing high-end semiconductor production – despite sustained public investment – suggested that restrictions on advanced chips and fabrication equipment could meaningfully slow China's progress in AI development. Early US policy measures under the first administration of President Donald Trump targeted individual firms, reflecting uncertainty about the durability and scope of semiconductor leverage. Over time, however, policymakers increasingly concluded that control over the infrastructure supporting AI development, rather than individual technology transfers, would determine long-term competitive advantage.

This logic culminated in President Joe Biden administration's October 2022 export controls, which restricted China's access to advanced AI chips and the equipment required to manufacture them. Subsequent coordination with allies expanded these restrictions to include lithography systems, semiconductor manufacturing tools, and advanced packaging technologies. Rather than functioning as traditional arms control, these measures sought to influence the pace, scale, and direction of Chinese technological development by limiting access to the physical infrastructure required to sustain AI training and deployment. In this sense, semiconductor export controls operate as a form of indirect technological governance, shaping rival states' development pathways rather than prohibiting specific applications outright.

Yet the attempt to weaponize semiconductor chokepoints reveals important limits to structural technological power. Unlike nuclear or conventional military technologies, advanced AI hardware circulates through dense commercial ecosystems characterized by multinational supply chains, private-sector innovation, and dual-use applications. These characteristics make technological denial difficult to sustain and generate significant tradeoffs between security objectives and economic competitiveness. Restrictions on semiconductor exports can impose costs on domestic firms, incentivize supply-chain diversification by targeted states, and complicate alliance coordination when partner countries bear disproportionate economic burdens.

Policy debates within the United States reflect these tensions. The Biden administration's export control regime emphasized technological denial and long-term strategic constraint. By contrast, the Trump administration's subsequent decision to permit sales of high-performance Nvidia H200-class chips to Chinese customers reflected a more permissive approach that prioritized maintaining US commercial influence within global AI ecosystems. Advocates of partial export liberalization argue that allowing Chinese firms access to “good enough” computing

infrastructure generates revenue for US companies, reinforces dependence on American software and programming platforms, and reduces incentives for China to accelerate indigenous semiconductor development. Critics counter that such policies risk narrowing the United States' compute advantage and accelerating Chinese AI capability development.

DeepSeek demonstrates that Chinese firms remain capable of producing highly competitive AI systems even under partial infrastructure constraints. While such models do not consistently match the performance of leading US frontier systems, they highlight how technological restrictions may redirect innovation toward efficiency, cost optimization, and alternative training strategies rather than preventing capability development altogether. These dynamics reinforce the reality that chokepoint power rarely produces absolute technological denial and instead operates by shaping competitive timelines and development pathways.

Semiconductor statecraft also carries broader geopolitical consequences. Export controls and technology restrictions send signals that rival states must interpret under conditions of uncertainty. Measures intended to slow technological progress may be read simultaneously as temporary bargaining leverage, long-term containment strategy, or acknowledgment of eventual convergence. This ambiguity complicates strategic signaling and can intensify reciprocal investment in domestic technological capabilities, reinforcing broader patterns of technological competition.

Finally, the effectiveness of semiconductor leverage depends on factors that extend beyond supply-chain control alone. Across both military and commercial technological revolutions, early pioneers have frequently struggled to translate technological breakthroughs into sustained institutional advantage. Leadership in artificial intelligence may therefore depend less on exclusive access to hardware or models than on the capacity to absorb AI systems across military organizations, commercial industries, and public-sector institutions. As scholarship on general-purpose technologies suggests, long-term competitive advantage often accrues to actors that successfully integrate new technologies into economic production, governance systems, and operational practice rather than those that simply pioneer technological breakthroughs.

Domestic Political Economy and Infrastructure Mobilization

Although the United States and China confront many of the same material constraints at the AI frontier, their ability to manage and mobilize the underlying infrastructure differs sharply. These differences reflect broader features of domestic political economy and state capacity that condition how each country can pursue control over AI-related chokepoints.

China's strategy for developing artificial intelligence involves coordinated policy and investment initiatives that extend beyond research firms to include the broader infrastructure supporting generative AI. In 2024, China's generative AI ecosystem was expanding rapidly, with more than

fifty domestic developers of large language models alongside major technology platforms and venture capital funding. They reflected a national emphasis on scaling both technological capabilities and the underlying compute infrastructure needed to support them (Triolo and Schaefer 2024). Unlike markets where infrastructure expansion is driven mainly by private hyperscalers, Chinese authorities have actively shaped the environment through regulatory approvals and institutional engagement that facilitate access to computing capacity, model deployment, and enterprise adoption. This state-aligned environment enables Chinese firms to iterate and deploy AI models while navigating constraints such as export controls on advanced chips. Although China's ecosystem remains diffuse compared with Western counterparts, the combination of public policy support, broad participation by domestic tech champions, and centralized regulatory oversight helps align investment and infrastructure development with national goals for technological competitiveness and economic integration.

The United States faces a different set of structural constraints in scaling AI infrastructure. Although US firms lead in model development, the physical expansion of data centers is shaped by decentralized regulatory and energy systems. Data-center siting is governed by local zoning processes, environmental reviews, water access constraints, and fragmented regional electricity markets that complicate large-scale infrastructure planning. Empirical analyses of US data-center energy consumption indicate that rapid growth in compute demand is already placing increasing strain on regional power systems, requiring substantial upgrades to transmission capacity and generation resources to sustain continued expansion (Shehabi et al. 2024). More broadly, recent policy scholarship argues that the US model of privately driven AI innovation has reached structural limits, as infrastructure expansion, energy provisioning, and regulatory coordination increasingly require state involvement to sustain continued scaling (Buchanan and Collins 2025).

These energy requirements intensify local political contestation, as municipalities, environmental groups, and state regulators weigh economic benefits against concerns over electricity costs, water usage, and grid reliability. Although local governments are often drawn to data-center projects by the prospect of tax revenue and construction employment, research suggests that traditional data-center development models frequently produce limited long-term job creation and uneven regional economic spillovers. This increases scrutiny over whether communities absorb infrastructure and environmental costs without sustained local benefits.

Rising demand for large-scale AI facilities has intensified competition for access to power, permits, and developable land, giving local and regional governments greater leverage in negotiations with technology firms over workforce investments, community benefit agreements, and infrastructure upgrades. US federal authority over land use and electricity markets remains limited. So, these localized bargaining dynamics introduce friction that national industrial or technology policy cannot easily override. This creates the possibility that the United States could face bottlenecks in compute infrastructure deployment even while maintaining leadership in AI model innovation (Goetzl, Muro, and Methkuppally 2026).

Those local debates have percolated into national-level political discussions, where elected leaders ranging from US Senator Bernie Sanders to Florida Governor Ron DeSantis, from different ends of the political spectrum, have called for either moratoria on AI data centers or expanded local authority to block their construction (Kimball 2026). Because data centers constitute the physical backbone of large-scale AI training and deployment, such domestic political conflicts can shape the scale and speed at which compute capacity expands. In turn, limits on compute expansion influence semiconductor demand, cloud infrastructure investment, and the broader technological ecosystems that underpin US–China strategic competition. Semiconductors have therefore shifted from a largely technical concern to a focal point of US–China strategic interaction that is also increasingly salient in domestic US politics. These debates illustrate how domestic political constraints may shape not only the pace of AI development but the United States’ capacity to sustain leadership in a globally competitive and infrastructure-intensive technological order.

Non-Great Powers

Competition over AI capabilities is increasingly concentrated among great powers – particularly the United States and China – whose advantages rest on control over compute infrastructure, semiconductor supply chains, and large-scale model development. Yet dominance in these structural foundations does not translate into monopoly over military AI innovation. Many of the most visible battlefield applications of AI have emerged in conflicts involving middle powers and regional actors. This pattern reflects a broader dynamic in which military AI diffuses operationally across the international system even as frontier research and infrastructure remain concentrated among technologically advanced states.

This layered diffusion occurs because military AI capabilities often depend less on frontier model development than on the integration of commercially available technologies, rapid doctrinal adaptation, and experimentation under combat conditions. While great powers dominate the upstream components of AI production, operational innovation frequently emerges among militaries actively engaged in conflict, where incentives to adopt emerging technologies, modify tactics, and accept operational risk are especially strong. Battlefield diffusion therefore demonstrates that structural technological leadership does not guarantee battlefield dominance, particularly in the early stages of technological transformation.

For advanced economies with relatively small populations – such as Israel, Singapore, Germany, and South Korea – AI-enabled systems offer a means of weakening the traditional link between manpower and military effectiveness. By automating labor-intensive functions including intelligence analysis, surveillance, logistics coordination, cyber operations, and elements of precision strike, these states can substitute capital, data, and organizational integration for large standing forces. Rather than competing at the technological frontier in model development, such states may derive advantage through rapid adoption, doctrinal flexibility, and effective

integration of commercially available systems (Borchert, Schütz, and Verbovszky 2024). In this sense, AI diffusion is reshaping military effectiveness horizontally across the international system, reducing the premium on large national populations for certain military missions even as frontier advantages remain concentrated among major powers (Raska and Bitzinger 2023; Sweijs, van Genugten, and Osinga 2024).

Israel's use of AI decision-support tools during the war in Gaza illustrates the operational appeal and controversy surrounding such systems. Designed to generate targeting recommendations at speeds beyond human analytic capacity, Israel's Lavender system reportedly provided significant operational efficiencies. Some observers argue that the system's speed and scale raised concerns about compliance with international humanitarian law, particularly when human operators lacked visibility into how targeting recommendations were generated (Dorsey 2024). Other legal assessments emphasize that Lavender functioned as decision-support rather than autonomous targeting, with human operators retaining authority over strike decisions (Schmitt 2024). Regardless of these legal debates, the case highlights how AI systems can transform battlefield decision tempo even when human oversight remains formally intact.

The Russia–Ukraine war further demonstrates how states operating under resource and technological constraints can rapidly innovate in AI-enabled military systems (Chivers 2025). Ukraine has emphasized low-cost adaptation, integrating algorithmic tools for reconnaissance, targeting assistance, and drone coordination to offset conventional force disadvantages. Russia has pursued a more experimental approach, testing AI-enabled autonomy in unmanned ground vehicles and deploying algorithmic decision-support tools in cyber and electronic warfare operations, effectively using the conflict as a live laboratory for evaluating AI's battlefield utility (Konaev 2023). These cases illustrate how active conflict accelerates AI experimentation in ways that peacetime procurement and testing processes often constrain.

Despite this rapid operational innovation, most drone systems employed by both sides remain remotely piloted or semi-autonomous rather than fully autonomous weapons. First-person-view (FPV) drones function primarily as inexpensive remote sensing and strike platforms in which human operators retain real-time control. Their effectiveness derives less from advanced autonomy than from low cost, mass deployment, and persistent surveillance. For example, during Operation Spider's Web, Ukrainian forces reportedly used image-recognition algorithms trained on datasets of Russian bombers to assist targeting, while human operators maintained direct strike control (Horowitz, Kahn, and Schwartz 2025). More broadly, remotely piloted one-way attack drones have accounted for an estimated 60–70 percent of frontline casualties in the conflict (Santora et al. 2025; Chivers 2024). However, continued electronic warfare and communications disruption have begun pushing both sides toward drones capable of greater onboard autonomy, suggesting a gradual shift toward more independent lethal systems when communication links degrade (Chivers 2025).

Battlefield diffusion also extends beyond active warfighting into surveillance and proxy warfare. India has incorporated AI into border monitoring along the Line of Actual Control with China, deploying automated image-recognition systems as part of its broader “AI in Defence” roadmap (Indian Ministry of Defence 2023; Bommakanti et al. 2024). Across the Middle East, states such as the United Arab Emirates and Saudi Arabia have imported Chinese AI-enabled drone platforms, including the Wing Loong II and Rainbow CH-4, which have been deployed in proxy conflicts in Libya and Yemen (Kreps and Rogers 2025). These cases illustrate how AI capabilities can spread through arms exports, security partnerships, and commercial technology transfers, further accelerating operational diffusion beyond the states leading AI research.

These diffusion dynamics extend beyond state militaries. Non-state armed groups have demonstrated increasing ability to adapt commercially available AI-enabled tools, particularly in drone warfare, surveillance, and information operations. The relatively low cost of commercial drones, combined with open-source computer vision and navigation software, allows militant organizations and proxy forces to deploy capabilities that previously required state-level research and procurement infrastructure. Policy analysis suggests that rapidly improving and widely accessible AI models can generate “malicious uplift” by amplifying the scale and speed of harmful activities and enabling capabilities previously limited to advanced military or intelligence organizations to become available to smaller groups or even individuals. Such tools can accelerate operational planning, enhance targeting and surveillance capabilities, and expand information operations and cyber activities across a broader range of actors. While these groups rarely operate at the frontier of AI development, their ability to rapidly adopt and adapt commercially available technologies complicates deterrence, attribution, and escalation management by expanding the range of actors capable of deploying AI-enabled military capabilities (Chan et al. 2026).

Taken together, these cases demonstrate that military AI adoption follows multiple diffusion pathways, including organizational integration, wartime experimentation, export-driven proliferation, and surveillance-based deployment. This diffusion complicates efforts to assess relative military capability because states may achieve operational advantages through integration and experimentation even when they lack leadership in frontier model development or compute infrastructure. Battlefield diffusion therefore reinforces the broader challenge of measuring technological advantage in AI and underscores the growing disconnect between structural technological leadership and operational military effectiveness.

Information and Cyber Superiority in an AI Era

Artificial intelligence also intersects with international security through its effects on information power. As AI is arguably the most important “new” technology of the last generation, how cyber interacts with AI is therefore critical to understand. AI reshapes how states collect, process, and act on data, and it alters the conduct of cyber operations that depend on digital infrastructure.

This section examines two related but distinct questions: How does AI affect the balance of information power in intelligence and military decision-making? And how does AI influence the offense–defense balance in cyberspace.

In this chapter, cyberspace refers not merely to the public internet but also the broader digital ecosystem of networked information systems, cloud infrastructure, software platforms, data storage architectures, and communications networks. AI both depends on and transforms this environment. The same digital systems that enable AI-driven analysis also create new vulnerabilities to intrusion, manipulation, and deception.

AI and the balance of information power

Artificial intelligence is transforming the balance of information power by expanding the speed, scale, and scope through which states collect, process, and interpret data. Since the end of the Cold War – particularly following the 1991 Persian Gulf War – information superiority has been viewed as a defining element of military and national power. AI intensifies this trend by enabling intelligence and military organizations to manage data volumes that exceed human analytic capacity.

AI-enabled analytic systems expand analysts’ ability to process large volumes of unstructured data, automate translation and transcription, identify behavioral patterns, and detect emerging narratives across digital information environments (Horowitz and Greenberg 2022; Goldstein et al. 2023). These tools can also integrate information across bureaucratic stovepipes, allowing intelligence organizations to connect disparate data streams more rapidly and potentially reduce the risk of strategic surprise (Zegart 2007).

Early operational experimentation illustrates how these capabilities are beginning to appear in national security contexts. During U.S. and allied operations related to Iranian-backed militia activity in the Middle East, reports indicated that large language models derived from commercial frontier systems, including models related to Anthropic’s Claude, were used experimentally to assist with the synthesis of large volumes of intelligence reporting, open-source information, and battlefield updates. Rather than generating targeting decisions directly, these systems helped analysts summarize reporting streams, identify relevant entities, and highlight potential patterns across large data environments that would otherwise require substantial analytic labor. Such applications illustrate how commercial frontier models can function as analytic infrastructure within military and intelligence workflows, accelerating information processing while leaving final judgments to human operators.

AI’s significance in intelligence, therefore, lies less in revolutionary front-end breakthroughs than in its ability to restructure the analytic infrastructure that supports decision making. AI compresses analytic timelines, filters informational relevance, and standardizes interpretive steps

that were previously manual, uneven, and labor intensive. These technologies do not replace human judgment, but they fundamentally condition it by shaping how information enters analytic pipelines, how quickly it moves through bureaucratic systems, and how it is implicitly prioritized.

Military organizations are increasingly integrating these capabilities into operational planning and battlefield decision making. The U.S. Department of Defense's Raven Sentry pilot program, for example, used AI trained on open-source data and historical insurgent attack patterns to predict likely insurgent activity in Afghanistan (Spahr 2024). More recent initiatives such as the Army Intelligence Data Platform seek to fuse intelligence from multiple sources and distribute it more rapidly to tactical commanders. Field experiments conducted by the Army's 75th Innovation Command and the 18th Airborne Corps, including exercises under the Maven Smart System and Scarlet Dragon programs, explore how AI tools can fuse sensor data, operate under degraded communication environments, reduce cognitive load on analysts, and support lethal and non-lethal targeting decisions.

Reflecting these developments, many policymakers and intelligence scholars argue that states capable of integrating AI across intelligence and operational decision-making processes may gain decisive informational advantages. Amy Zegart, for example, argues that AI offers potentially transformative capabilities by allowing organizations to process data at unprecedented scale, identify patterns more rapidly, and augment human analytic performance (Zegart 2022). Some observers have suggested that these technologies may even reduce or prevent intelligence failures by enabling faster detection of adversary intentions (Moran et al. 2023).

The Vulnerability Paradox of AI-Enabled Information Power

AI-enabled information systems introduce new forms of fragility. Advanced models often function as opaque computational processes, making it difficult for analysts to interpret or audit machine-generated conclusions. They can also be brittle, performing reliably under familiar conditions while failing unpredictably when confronted with adversarial or novel inputs (Goldfarb and Lindsay 2021).

More significantly, AI systems create new attack surfaces within intelligence and military organizations. Adversaries may attempt to corrupt training datasets, manipulate model inputs, or exploit architectural weaknesses through cyber operations. Because AI systems increasingly operate within complex "systems of systems," malicious actors may target any component of the analytic pipeline. Machine-generated assessments may therefore become vulnerable to adversarial deception designed to manipulate decision-makers indirectly (Maschmeyer 2024).

These dynamics produce a paradox of AI-driven information power. As states rely more heavily on automated analytic systems to enhance situational awareness and decision speed, they may simultaneously increase exposure to deception strategies designed to exploit those systems. Informational advantage and informational vulnerability expand together.

AI and the Cyber Offense–Defense Balance

Artificial intelligence also reshapes competition in cyberspace itself. Here, the relevant question is whether AI shifts the offense–defense balance in digital conflict. Traditional offense–defense theory holds that when offensive capabilities are dominant, conflict becomes more likely because conquest appears easier (Jervis 1978; Van Evera 1998). This framework has long been debated, particularly in cyber scholarship. Early analyses frequently described cyberspace as inherently offense-dominant due to speed, low barriers to entry, and ambiguity between espionage and attack (Buchanan 2016). More recent research challenges this view, showing that sophisticated offensive cyber campaigns require organizational capacity, skilled personnel, and sustained operational secrecy (Slayton 2016; Lonergan and Lonergan 2023; Smeets 2022). Cyber operations often succeed through deception and intelligence preparation rather than brute technical capability (Gartzke and Lindsay 2015).

AI may alter this balance in complex ways. Some analysts argue that AI strengthens cyber offense by accelerating attack planning, lowering technical barriers, and increasing the scale of operations. Machine learning can automate reconnaissance, generate adaptive malware, and speed exploitation processes that previously required manual intervention (Bonfanti 2022; Lohn 2025). A RAND analysis warned that AI-enabled cyber capabilities could enable a “splendid first cyber strike” capable of disrupting retaliatory responses (Mitre and Predd 2025).

At the same time, AI may enhance cyber defense. AI-enabled cybersecurity tools improve network monitoring, automate vulnerability detection and patching, and strengthen anomaly detection. Because defenders control network architecture, AI may allow them to dynamically reshape digital terrain and increase the cost of intrusion. Some scholars argue that AI may offer near-term advantages to defenders capable of integrating automated security systems effectively (Garfinkel and Dafoe 2019).

Empirical evidence suggests a mixed pattern. State-sponsored cyber actors increasingly use AI for reconnaissance and social engineering, yet human oversight remains central to strategic decision making. AI appears to increase automation and efficiency rather than eliminate the need for human control. As in the intelligence domain, AI reshapes competition incrementally rather than producing decisive shifts.

AI Governance

Structural Constraints on AI Governance

Given its enormous implications for international politics, the spread of AI – especially frontier models and military applications – forces states to confront whether and how AI can be regulated, constrained, or stabilized through international mechanisms. Yet AI poses governance challenges that differ fundamentally from those associated with earlier arms control efforts. Unlike nuclear weapons, which depended on rare materials, centralized infrastructure, and observable tests, AI lacks chokepoints that reliably support verification and enforcement of binding international agreements. While AI development depends on upstream supply chains such as advanced semiconductors and compute infrastructure that can serve as sites of strategic competition and export control, these chokepoints are deeply embedded in global commercial ecosystems and dual-use markets, making them far less suitable anchors for traditional arms control regimes.

Patterns of diffusion further complicate governance by shaping how states interpret one another's capabilities and intentions. Some AI models remain open-source and freely accessible, while others operate within closed, proprietary ecosystems that require paid access or state cooperation. Chinese developers have often released leading models that are open-source (Meinhardt 2025). By contrast, many leading U.S. frontier models remain proprietary, reflecting a model of concentrated control within private firms (Eastwood 2026). Open models accelerate diffusion, allowing states, firms, and non-state actors to reproduce capabilities quickly, but that same openness frustrates control, monitoring, and oversight. Closed systems, by contrast, concentrate advantage among a small number of firms and states and slow diffusion, yet they also restrict transparency and leave rivals uncertain about what those systems can do. Both pathways heighten mistrust in different ways: Openness raises fears of rapid proliferation, while closed development fuels suspicion of concealed advances.

As Jane Vaynman and Tristan Volpe (2025) argue, these factors combine to create an “arms control dead zone,” where the tools of past regimes – definitions, inspections, taboos, treaties – have no clear anchor. AI governance thus begins from fragmentation rather than consensus, with states disagreeing not only about acceptable uses of AI but about where risks originate, whether from widespread access, concentrated control, or organizational misuse.

Not surprisingly, then, the governance landscape remains fragmented and has struggled to keep pace with the rate of technological change. AI lacks binding treaties, strong taboos, or shared verification mechanisms to date. States disagree not only about acceptable uses of AI but also about where risks originate – whether from widespread access, concentrated control, or organizational misuse. As a result, governance initiatives rarely begin from shared premises.

Instead of pursuing comprehensive regulation of AI as a technology, states and institutions advance narrower, partial approaches that reflect their own strategic priorities, regulatory traditions, and threat perceptions.

Civilian / General AI Governance: Partial and Indirect Approaches

In response to these constraints, international AI governance has converged around a limited set of strategies, each attempting to impose structure on a technology that resists clear boundaries. One starting point has been scope and definition. Rather than regulate “AI” writ large, many proposals seek to delimit categories of particular concern. Some draw thresholds around model size or training compute, as in the EU AI Act. It imposes stricter obligations on general-purpose models exceeding 10^{23} FLOPs (EU AI Act 2024), an approach similar to that taken by a Biden executive order that was later overridden by the Trump administration.

The difficulty is structural rather than technical. No widely accepted, objective metric currently captures “increased capability” across heterogeneous domains such as reasoning, autonomy, cyber exploitation, or biological knowledge. Static FLOPs cutoffs therefore function as administrable proxies, not definitive indicators of risk. As techniques evolve, fixed numerical thresholds risk becoming either overinclusive or obsolete, requiring continual revision rather than providing durable regulatory clarity.

A second cluster of proposals shifts attention from definitions to verification and transparency. Because training frontier systems requires substantial computational resources, some analysts have suggested monitoring the movement and use of advanced chips, analogizing compute to fissile material. Proposals include registries of large training runs, mandatory reporting requirements, and third-party audits of training processes or safety evaluations (Dafoe et al. 2023). In addition, some policy discussions have explored whether hardware-level mechanisms – such as embedded usage reporting, geofencing constraints, or remote management features built into advanced AI chips – could provide governments with greater visibility into how high-end compute is deployed. In principle, such mechanisms might allow exporting states to monitor or restrict the use of frontier chips abroad.

The goal of these approaches is not comprehensive enforcement but sufficient visibility to reduce mistrust and deter reckless behavior. Yet structural obstacles loom large. Unlike nuclear materials, advanced chips circulate through global commercial markets; firmware and system architectures can be modified; and firms may resist embedding intrusive monitoring features that threaten intellectual property or expose them to geopolitical retaliation. As RAND analysts observe (2023), nuclear verification succeeded because of identifiable chokepoints; AI supply chains, by contrast, are embedded in dense and highly distributed commercial ecosystems, making hardware-based transparency technically complex and politically fraught.

Recognizing these limits, a third set of proposals turns to new institutions and indirect levers. Some scholars and policymakers have advocated creating an international AI agency modeled on the International Atomic Energy Agency (IAEA) to monitor frontier models and mediate disputes (Belfield 2025). It is an idea echoed by the UK government at the 2023 Bletchley Park AI Safety Summit, through calls for a global AI Safety Institute network. Critics, however, argue that such an institution would struggle for the same reasons outlined above: As a general-purpose technology embedded in civilian markets, AI cannot be monitored or controlled in the way nuclear programs can (Horowitz and Kahn 2025).

In practice, governance has increasingly shifted toward indirect and hybrid mechanisms. National policy instruments now function as de facto global governance tools. Export controls on advanced semiconductors – discussed earlier in the context of US–China competition – shape the pace and direction of AI development by constraining access to large-scale training infrastructure (Miller 2022). The Biden administration’s October 2022 chip rules illustrate how licensing regimes tied to compute intensity can operate as forms of indirect control even without formal international agreement. Alongside these state-centered measures, governments have pursued multistakeholder approaches aimed at managing risk where binding regulation is infeasible.

In 2023, the White House secured voluntary commitments from leading AI firms on model evaluation, information sharing, and deployment practices. At the industry level, frontier developers themselves have moved to institutionalize coordination through the Frontier Model Forum, which brings together leading platform providers to develop shared safety practices, technical benchmarks, and approaches to risk assessment for the most advanced models.

Internationally, the Frontier AI Safety Forum launched after the Bletchley Park Summit brought together governments and frontier developers to coordinate research, share risk assessments, and develop common safety benchmarks. The UN High-Level Advisory Body on AI has similarly emphasized voluntary standards and coordination over legal compulsion.

Underpinning these initiatives is a growing reliance on norms and standards as substitutes for formal regulation. Technical benchmarks, safety evaluations, reporting practices, and certification schemes do not constrain capability development directly, but they shape expectations about responsible behavior and acceptable risk. Over time, such standards can harden into de facto requirements through procurement rules, liability regimes, and market access, creating forms of governance that operate without treaty obligations or centralized enforcement.

Across these efforts, standards and norms increasingly serve as governance substitutes where treaties cannot. Trager et al. (2023) argue that licensing, certification, and oversight standards can function as a global scaffold for civilian AI governance by creating incentives for compliance through market access and reputational effects. Rather than restraining AI development directly, such mechanisms seek to shape behavior at key points of deployment and

diffusion. Horowitz and Kahn (2025) argue that standards-setting approaches used for general-purpose technologies, such as those developed through the International Telecommunications Union, are more likely to succeed in the AI context than traditional arms control treaties.

The effectiveness of standards-based governance also depends heavily on sustained political commitment and international credibility. Periods of domestic policy volatility, shifting regulatory priorities, or skepticism toward multilateral coordination may complicate efforts to build durable global norms, while also creating space for competing governance models to emerge. As a result, standards-based approaches may represent a promising but politically contingent pathway for managing AI competition.

Taken together, these early efforts to apply international relations theories to the question of AI governance suggest that attempts at comprehensive regulation could founder on problems of definition, verification, and enforcement, as well as governance shifts toward leverage, coordination, and norm formation. This evolution reflects not a lack of ambition but an accommodation to the structural realities of AI as a general-purpose, dual-use technology embedded in global markets. Governance, where it occurs, is therefore partial, indirect, and contested – aimed at shaping trajectories and managing risk at the margins rather than controlling or banning, which have been out of reach. An interesting question for future researchers will be identifying the conditions in which these limits to governance approaches could change.

Military AI Governance: Application-Specific and Normative Efforts

If governance of civilian AI is fragmented and indirect, governance of military AI is even more constrained, though attempts at binding international legal regulation in specific areas are more advanced. Statutory regimes such as the EU AI Act largely sidestep military and national security uses, leaving governance efforts to concentrate on specific applications and norms rather than comprehensive control. Existing approaches to military AI governance fall into three broad categories. The first centers on use-case-specific regulation, most prominently the discussions on lethal autonomous weapons systems (LAWS) within the UN Convention on Certain Conventional Weapons (CCW). The LAWS Group of Governmental Experts (GGE) represents an effort to isolate a particular application of AI and assess its compatibility with international humanitarian law (United Nations Office of Disarmament Affairs 2025). While these talks have clarified areas of concern, they have also revealed deep disagreement over definitions, scope, and the feasibility of meaningful restrictions, limiting progress beyond general principles. For example, the land mines and cluster munitions treaty efforts proceeded by agreement on a ban and then efforts to define specifically what was being banned. However, in a world where LAWS could encompass anything from loitering munitions to autonomous tanks, that approach has

proven less successful as states have sought more clarity earlier in the process on what, specifically, would be regulated.

It is possible that these LAWS governance efforts could evolve in 2026, given the call by the UN Secretary General for binding regulations on LAWS use. Specifically, civil society groups are advocating for a binding treaty, whether through the United Nations or an ad hoc process such as the Treaty on the Prohibition of Nuclear Weapons (TPNW). Consistent with international relations theories about civil society and transnational activism, if the CCW discussions continue to be slower than advocacy groups would like, civil society groups will likely attempt to rally supportive countries to convene an ad hoc process and negotiate a treaty to ban at least some uses of LAWS (Price 2003, Ritchie and Kmentt 2021).

A second category encompasses broader efforts to articulate norms and expectations for military AI use without binding legal commitments. The Responsible AI in the Military Domain (REAIM) process exemplifies this approach. Through a series of summits, REAIM has brought together states, international organizations, and civil society to emphasize human responsibility, legal compliance, and risk awareness in military AI deployment. Related initiatives include the US-led Political Declaration on Responsible Military Use of AI and Autonomy, which has attracted more than sixty signatories, and the November 2024 US–China agreement to maintain human control over nuclear weapons. These efforts do not impose constraints on capability development, but they seek to shape how AI is integrated into military decision making before practices become entrenched (US Department of State 2024, Renshaw and Hunnicutt 2024).

The third category focuses less on restraint and more on coordination among aligned states. Alliances and informal groupings such as NATO and the Quad have increasingly treated AI as an operational priority, emphasizing interoperability, shared standards, and collective learning. NATO has issued principles on responsible AI use, while the Quad’s Critical and Emerging Technology Working Group has articulated shared commitments on AI development and deployment. These arrangements are not designed to limit AI use, but to ensure that allied militaries can operate together despite divergent national approaches (The Quad 2023, NATO 2021).

Across all three categories, a common constraint emerges. Unlike nuclear or chemical weapons, AI’s general-purpose character and wide diffusion make comprehensive, legally binding agreements hard to negotiate and even harder to enforce. A broad range of actors is adopting military AI, from major powers to middle states and non-state groups, each setting its own thresholds for acceptable use. As a result, states are experimenting in parallel rather than converging on binding shared rules. That said, they often are informally converging on shared norms for AI capability development, i.e. the importance of testing and evaluation to ensure that capabilities work effectively. Even among allies, differences in confidence levels, testing standards, and targeting practices have already created interoperability challenges.

From an international relations perspective, this pattern suggests that traditional arms-control models offer limited guidance. Instead, confidence-building measures, transparency practices, and shared testing norms appear to be the most plausible building blocks for cooperation (Horowitz and Kahn 2021). Proposals include exchanges of information on military AI programs, joint exercises designed to probe failure modes, incident hotlines, and agreements governing interactions between deployed autonomous systems, modeled loosely on Cold War-era incidents-at-sea arrangements (CNAS 2023; GC REAIM 2025). These measures aim not to eliminate competition but to reduce the risks of accident, misinterpretation, and unintended escalation.

Efforts to govern military AI reflect a shift away from control toward management under conditions of diffusion and uncertainty. The challenge is no longer bilateral, nor confined to weapons counts or deployment thresholds but one shaped by uneven adoption, organizational capacity, and opaque integration into command systems.

Conclusion

Artificial intelligence is reshaping international security not by introducing a discrete new instrument of force but by altering how power is generated, perceived, and exercised across economic, military, and informational domains. As a general-purpose technology embedded in civilian infrastructure, AI resists traditional categories of arms races and weapons competition. Its strategic significance lies less in measurable model performance than in how capabilities diffuse and are embedded across institutions. Across domains – from battlefield decision making to cyber operations and great-power rivalry – AI generates persistent ambiguity about relative advantage. States therefore compete and hedge even in the absence of clear shifts in material balance, amplifying familiar security-dilemma pressures without producing decisive dominance.

The analysis of military adoption highlights this point. AI-enabled systems are being incorporated across intelligence, decision support, surveillance, and semi-autonomous weapons, but in ways that are highly context dependent. States differ not only in access to models or compute but in organizational capacity, doctrine, and tolerance for risk. These differences shape how AI is actually used and, therefore, how much advantage it confers. The chapter's discussion of battlefield integration and cyber operations suggests that AI's effects are mediated through institutions rather than determined by technology alone. In both domains, optimism about speed and automation is tempered by vulnerabilities to deception, brittleness, and organizational mismatch. The result is not a stable shift toward offense or defense, but a contingent balance shaped by how technologies are embedded in practice.

The examination of US–China competition highlights a parallel dynamic at the geopolitical level. Leadership in AI increasingly depends on access to compute and the infrastructure required to deploy it at scale, anchoring competition in semiconductor supply chains, energy

systems, and industrial policy. These material foundations create chokepoints that offer leverage but not control. Export controls and allied coordination can slow and redirect development, but they operate indirectly and incompletely within global markets. At the same time, domestic political economy conditions shape each country's capacity to mobilize infrastructure, suggesting that long-run advantage will depend as much on governance and implementation as on innovation itself.

These features have important implications for AI governance. Traditional arms-control models, which rely on clear definitions, centralized production, and observable verification, offer limited guidance for a technology that is diffuse, rapidly evolving, and embedded in civilian systems. As a result, governance has shifted toward partial and indirect mechanisms: export controls, standards, voluntary commitments, and coordination among firms and allies. In the military domain, efforts have focused on application-specific regulation, norm articulation, and confidence-building measures rather than comprehensive prohibition. Across both civilian and military contexts, governance proceeds through fragmentation rather than consensus, reflecting disagreement not only about acceptable uses of AI but about where risks originate and how they should be managed.

AI competition is unlikely to yield clear equilibria or decisive technological victories. Instead, it will remain characterized by uneven adoption, contested signals, and ongoing uncertainty about capability, intent, and restraint. Advantage will accrue not simply to those who innovate first or control key inputs, but to those who can integrate AI systems effectively across institutions, manage their vulnerabilities, and interpret rivals' actions under conditions of limited information.

For scholars of international relations, this implies that AI should be understood less as a singular driver of change than as a force multiplier for existing structural and organizational dynamics. For policymakers, it suggests that efforts to govern AI will be incremental, contested, and shaped by trade-offs rather than comprehensive control. And for future research, it highlights the need to move beyond technological benchmarks toward systematic study of adoption, integration, and institutional learning – where the strategic consequences of artificial intelligence are ultimately decided.

References

- Betts, Richard K. 1999. "Must War Find a Way?: A Review Essay." *International Security* 24(2): 166-198.
- Biddle, Stephen. 2006. *Military Power: Explaining Victory and Defeat in Modern Battle*. Princeton University Press.

- Bommakanti, Kartik; Yogesh Joshi; Shimona Mohan; Karthik Nachiappan; and Antara Vats. 2024. *AI in Defence and Border Monitoring: India's Use of Artificial Intelligence Along the Line of Actual Control*. Observer Research Foundation, May 10, 2024. <https://www.orfonline.org/public/uploads/posts/pdf/20240510151332.pdf>
- Bonfanti, Matteo E. 2022. "Artificial intelligence and the offence-defence balance in cyber security." *Cyber Security: Socio-Technological Uncertainty and Political Fragmentation*. London: Routledge: 64-79.
- Buchanan, Ben. 2016. *The cybersecurity dilemma: Hacking, trust, and fear between nations*. Oxford University Press.
- Buchanan, Ben, and Tantum Collins. 2025. "The AI Grand Bargain: What America Needs to Win the Innovation Race." *Foreign Affairs*, November/December 2025.
- Buzan, Barry, and Eric Herring. 1998. *The Arms Dynamic in World Politics*. Boulder, CO: Lynne Rienner Publishers.
- Campaign to Stop Killer Robots. 2023. *The Case for a Preemptive Ban on Fully Autonomous Weapons*. Accessed September 6, 2025. <https://www.stopkillerrobots.org>.
- Chan, Kyle, Michael E. O'Hanlon, Qi Haotian, and Zheng Lefeng. 2026. *AI Risks from Non-State Actors*. Brookings Institution.
- Chivers, C.J. "How Suicide Drones Transformed the Front Lines in Ukraine," *The New York Times Magazine*, December 31, 2024, <https://www.nytimes.com/2024/12/31/magazine/drones-weapons-ukraine-war.html>.
- Chivers, C.J. "In Ukraine, a New Arsenal of Killer A.I. Drones is Being Born," *The New York Times*, December 31, 2025, <https://www.nytimes.com/2025/12/31/magazine/ukraine-ai-drones-war-russia.html>.
- CNAS (Center for a New American Security). 2023. *Artificial Intelligence and International Stability: Risks and Confidence-Building Measures*. Washington, DC: CNAS. <https://www.cnas.org/publications/reports/ai-and-international-stability-risks-and-confidence-building-measures>.
- Copp, T., Dwoskin, E., & Duncan, I. (2026, March 4). Anthropic's AI tool Claude central to U.S. campaign in Iran, amid a bitter feud. *The Washington Post*. <https://www.washingtonpost.com/technology/2026/03/04/anthropic-ai-iran-campaign>

Dafoe, Allan, et al. 2023. *Verification for International AI Governance*. International Alliance for Partnerships in AI (IAPS). https://www.iaps.ai/s/Verification_for_International_AI_Governance.pdf.

Dafoe, Allan. 2018. *AI Governance: A Research Agenda*. Oxford: Future of Humanity Institute.

Dorsey, Jessica. "Israel's AI-Enabled Targeting of Hamas Members Jeopardizes Moral and Legal Standards of Warfare." *Utrecht University*, 18 July 2024. <https://www.uu.nl/en/achtergrond/israels-ai-enabled-targeting-of-hamas-members-jeopardizes-moral-and-legal-standards-of-warfare>

Eastwood, Brian. "AI Open Models Have Benefits. So Why Aren't They More Widely Used?" *MIT Sloan Management Review*, 20 Jan. 2026.

European Union. 2024. *Regulation (EU) 2024/1689 of the European Parliament and of the Council on Artificial Intelligence (AI Act)*. Official Journal of the European Union, L. 275.

Farrell, Henry, and Abraham L. Newman. "Weaponized Interdependence: How Global Economic Networks Shape State Coercion." *International Security* 44, no. 1 (Summer 2019): 42–79.

Garfinkle, Ben and Allan Dafoe. 2019. "Artificial Intelligence, Foresight, and the Offense-Defense Balance." *War on the Rocks*. <https://warontherocks.com/2019/12/artificial-intelligence-foresight-and-the-offense-defense-balance/>.

Gartzke, Erik, and Jon R. Lindsay. 2015. "Weaving tangled webs: offense, defense, and deception in cyberspace." *Security Studies* 24(2): 316-348.

Goetzel, D., Muro, M., & Methkuppally, S. (2026). *Turning the data center boom into long-term, local prosperity*. Brookings Institution.

Goldfarb, Avi, and Jon R. Lindsay. 2021. "Prediction and judgment: Why artificial intelligence increases the importance of humans in war." *International Security* 46(3): 7-50.

Goldstein, J. A., Sastry, G., Musser, M., DiResta, R., Gentzel, M., & Sedova, K. (2023). *Generative language models and automated influence operations: Emerging threats and potential mitigations*. arXiv. <https://doi.org/10.48550/arXiv.2301.04246>

Goodman, Peter. 2025. "18,000 Reasons It's So Hard to Build a Chip Factory in America," *New York Times*, <https://www.nytimes.com/2025/12/04/business/tsmc-phoenix-fab.html>

Hendrycks, Dan, Eric Schmidt, and Alexandr Wang. 2025. *Superintelligence Strategy: Expert Version*. arXiv. <https://arxiv.org/abs/2503.05628>.

Horowitz, Michael C. 2010. *The Diffusion of Military Power: Causes and Consequences for International Politics*. Princeton, NJ: Princeton University Press.

Horowitz, Michael C. 2018. "Artificial Intelligence, International Competition, and the Balance of Power." *Texas National Security Review* 1 (3): 36–57. <https://tnsr.org/2018/05/artificial-intelligence-international-competition-and-the-balance-of-power/>.

Horowitz, Michael C., and Erik Lin-Greenberg. 2022. "Algorithms and influence artificial intelligence and crisis decision-making." *International Studies Quarterly* 66(4).

Johnson, James. "Inadvertent escalation in the age of intelligence machines: A new model for nuclear risk in the digital age." *European Journal of International Security* 7.3 (2022): 337-359.

Johnson, James. "'Catalytic nuclear war' in the age of artificial intelligence & autonomy: Emerging military technology and escalation risk between nuclear-armed states." *Journal of Strategic Studies* (2021): 1-41.

Kaag, John and Sarah Kreps. *Drone Warfare*. 2014. Polity Press (New York, NY).

Kahl, Colin H. 2026. "The Myth of the AI Race: Neither America Nor China Can Achieve True Tech Dominance." *Foreign Affairs*, January 12, 2026.

Kimball, Spencer. "Bernie Sanders and Ron DeSantis Speak Out Against AI Data Centers Over Electricity Prices." *CNBC*, 1 Jan. 2026,

<https://www.cnn.com/2026/01/01/ai-data-centers-bernie-sanders-ron-desantis-electricity-prices.html>

Konaev, Margarita. *Use of AI and Autonomous Technologies in the War in Ukraine*. CNA, Oct. 2023,

www.cna.org/reports/2023/10/Use-of-AI-and-Autonomous-Technologies-in-the-War-in-Ukraine.pdf

Kreps, Sarah and James Patton Rogers. 2025. *Drones: What Everyone Needs to Know*. (Oxford, UK: Oxford University Press).

Jervis, Robert. 1978. "Cooperation under the Security Dilemma." *World Politics* 30 (2): 167–214. <https://doi.org/10.2307/2009958>.

Levy, Jack S. 1984. "The Offensive/Defensive Balance of Military Technology: A Theoretical and Historical Analysis." *International Studies Quarterly* 28: 219-238.

Lohn, Andrew J. 2025. "The Impact of AI on the Cyber Offense-Defense Balance and the Character of Cyber Conflict." arXiv. <https://arxiv.org/abs/2504.13371>.

Lonergan, Erica D. and Shawn W. Lonergan. 2023. *Escalation Dynamics in Cyberspace*. Oxford University Press.

Maschmeyer, Lennart. 2024. *Subversion: From Covert Operations to Cyber Conflict*. Oxford University Press.

Meinhardt, Caroline, Sabina Nong, Graham Webster, Tatsunori Hashimoto, and Christopher Manning. *Beyond DeepSeek: China's Diverse Open-Weight AI Ecosystem and Its Policy Implications*. Stanford Institute for Human-Centered Artificial Intelligence, 2025.

Ministry of Defence, Government of India. 2023. *AI in Defence: Roadmap for Adoption of Artificial Intelligence in the Indian Armed Forces*. New Delhi: Ministry of Defence. <https://www.ddpmod.gov.in/sites/default/files/2023-11/ai.pdf>

Mitre, Jim and Joel B. Predd. 2025. *Artificial General Intelligence's Five Hard National Security Problems*. RAND Corporation. https://www.rand.org/content/dam/rand/pubs/perspectives/PEA3600/PEA3691-4/RAND_PEA3691-4.pdf.

Moran, Christopher R, Joe Burton, George Christou. 2023. "The US Intelligence Community, Global Security, and AI: From Secret Intelligence to Smart Spying," *Journal of Global Security Studies*, 8(2).

National Intelligence Law of the People's Republic of China. Translated by ChinaLawTranslate.com, 2017, www.chinalawtranslate.com/en/national-intelligence-law-of-the-p-r-c-2017/. Accessed 22 Dec. 2025

NATO. 2021. "NATO Secretary General: Artificial Intelligence Will Be Key to Alliance Security." NATO Press Release, October 2021. https://www.nato.int/cps/en/natohq/news_187607.htm.

RAND Corporation. 2023. *Insights from Nuclear History for International AI Governance*. Santa Monica, CA: RAND. <https://www.rand.org/pubs/perspectives/PEA3652-1.html>.

Roff, Heather M., and Richard Moyes. 2016. *Meaningful Human Control, Artificial Intelligence and Autonomous Weapons*. London: Article 36. <https://article36.org/wp-content/uploads/2016/04/MHC-AI-and-AWS-FINAL.pdf>.

Santora, Marc, et al., “A Thousand Snipers in the Sky: The New War in Ukraine,” *The New York Times*, March 3, 2025, <https://www.nytimes.com/interactive/2025/03/03/world/europe/ukraine-russia-war-drones-deaths.html>.

Satariano, Adam and Paul Mozur. 2025. “Saudi Arabia’s New Power Play is Exporting A.I. to the World.” *The New York Times*, November 4, 2025. <https://www.nytimes.com/2025/10/27/technology/saudi-arabia-ai-exporter.html>.

Scharre, Paul. 2019. “Killer Apps: The Real Dangers of an AI Arms Race.” *Foreign Affairs*, April 16, 2019. <https://www.foreignaffairs.com/articles/2019-04-16/killer-apps>.

Scharre, Paul and Megan Lamberth. 2022. “Artificial Intelligence and Arms Control.” *arXiv*. <https://arxiv.org/abs/2211.00065>.

Schmitt, Michael N. 2024. “Israel – Hamas 2024 Symposium – The Gospel, Lavender, and the Law of Armed Conflict.” *Articles of War*, Lieber Institute for Law and Warfare, June 28, 2024. <https://lieber.westpoint.edu/gospel-lavender-law-armed-conflict/>

Shehabi, Arman, et al. *United States Data Center Energy Usage Report*. Lawrence Berkeley National Laboratory, 2024.

Slayton, Rebecca. 2016. "What is the cyber offense-defense balance? Conceptions, causes, and assessment." *International Security* 41(3): 72-109.

Smeets, Max. 2022. *No Shortcuts: Why States Struggle to Develop a Military Cyber-Force*. Oxford University Press.

Spahr, Thomas W. 2024. "Raven Sentry: Employing AI for indications and warnings in Afghanistan." *The US Army War College Quarterly: Parameters* 54(2).

State Council of the People’s Republic of China. *Made in China 2025*. Beijing, 2015.

Timbie, James, and Marjory S. Blumenthal. 2023. *Managing the Risks of Artificial Intelligence*. Washington, DC: Carnegie Endowment for International Peace.

Triolo, Paul, and Kendra Schaefer. “China’s Generative AI Ecosystem in 2024: Rising Investment and Expectations.” *National Bureau of Asian Research*, June 27, 2024.

The Quad. 2022. “Quad Principles on Technology Design, Development, Governance, and Use.” Critical and Emerging Technology Working Group, May 2022. <https://www.whitehouse.gov/briefing-room/statements-releases/2022/05/20/quad-principles-on-technology/>.

UK Government. 2023. *Bletchley Declaration on AI Safety*. November 1, 2023. <https://www.gov.uk/government/publications/bletchley-declaration>.

U.S. Department of War. *Artificial Intelligence Strategy for the Department of War*. Memorandum, January 9, 2026. Artificial Intelligence Strategy for the Department of War

UNIDIR (United Nations Institute for Disarmament Research). 2023. *Confidence-Building Measures for Military Artificial Intelligence*. Geneva: UNIDIR. <https://unidir.org/publication/confidence-building-measures-for-military-artificial-intelligence>.

Van Evera, Stephen. 1998. "Offense, defense, and the causes of war." *International Security* 22(4).

Vaynman, Jane, and Tristan Volpe. 2025. "Competition and Collusion: How the AI Arms Race Can Motivate Governance." In *RAND–Perry World House Collaboration Report*. RAND Corporation and Perry World House.

Udinmwen, Efosa, "Xi Jinping pushes China's AI ambition but warns against idle capacity," *Tech Radar*, January 30, 2026. <https://www.yahoo.com/news/articles/xi-jinping-pushes-chinas-ai-223500014.html>

White House. 2022. "Fact Sheet: Implementation of Export Controls on Advanced Computing and Semiconductor Manufacturing Items." October 7, 2022. <https://www.commerce.gov/news/fact-sheets/2022/10/fact-sheet-implementation-export-controls>.

United States. 2025. *U.S. Export Controls and China: Advanced Semiconductors*. Congressional Research Service Report (R48642), September 19, 2025. <https://www.congress.gov/crs-product/R48642>

White House. 2023. "Voluntary Commitments to Manage AI Risks." July 21, 2023. <https://www.whitehouse.gov/briefing-room/statements-releases/2023/07/21/fact-sheet-biden-harris-administration-secures-voluntary-ai-commitments/>.

White House. 2025. *America's AI Action Plan*. The White House, Executive Office of the President. <https://www.whitehouse.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf>

Zegart, Amy. 2007. *Spying Blind: The CIA, the FBI, and the Origins of 9/11*. Princeton University Press.

Zegart, Amy. 2022. *Spies, Lies and Algorithms: The History and Future of American Intelligence*. Princeton University Press.