

POLICY MEMO

Who Governs the Machine? Sovereignty, Power, and the Reordering of Authority in the Age of Frontier AI

TSIPORAH FRIED

Visiting Senior Fellow, Hudson Institute

July 2026

Since the emergence of large language models, artificial intelligence (AI) has consistently dominated headlines. Public debate has largely centered on its philosophical implications—its transformative impact on work, education, creativity, and democracy, as well as on what it means to be human. Yet in contrast to issues like cloud computing, semiconductor supply chains, digital infrastructure, data governance, or telecommunications networks, which policymakers have long recognized as matters of national security and geopolitical competition, AI was primarily understood as a societal revolution, one that raised profound ethical, economic, and existential questions.

Frontier AI's open architecture and broad accessibility, in addition to leading AI companies' universalist narrative, fostered the perception that foundation models would broadly benefit humanity rather than serve as instruments of

geopolitical competition. OpenAI, Anthropic, and other frontier laboratories consistently framed their mission in universal terms: building AI would serve humanity and transcend national boundaries, and the technology would ultimately require international governance commensurate with its global impact.

Developments in the past six months have fundamentally challenged that narrative. The confrontation between Anthropic and the United States Department of War (DoW) over the military use of Claude; the Trump administration's restrictions on access to Anthropic's Fable 5 and Mythos 5 models; OpenAI's acceptance of a government-supervised rollout of GPT-5.6; and the broader expansion of government oversight over the deployment of frontier models all point in the same direction. Even the Vatican's intervention through the papal encyclical *Magnifica Humanitas* reflects the recognition that AI

is more than a mere technological innovation. Instead, it raises questions about political authority, moral responsibility, and strategic power.

Because of these changes, governments cannot assume that frontier AI models are simply globally accessible digital services. Now, leaders are increasingly treating them as strategic assets whose development, deployment, and dissemination involve national security considerations and geopolitical competition, and are placing them under governmental oversight. The question has moved from whether frontier AI has strategic value, to who exercises authority over that value and according to which political and legal principles.

This marks the beginning of a new era in AI governance. The debate—once confined to innovation, safety, and regulation—has become a contest over sovereignty, legitimacy, strategic autonomy, and control of what may prove to be one of the defining strategic capabilities of the twenty-first century. Crucially, this competition is no longer taking place solely between states, specifically China and the United States. It is unfolding between governments and private companies that develop frontier AI models, each seeking to define the rules governing access to, deployment of, and control over these technologies.

A recent episode in June—when Washington shocked Europe by temporarily placing export controls on Anthropic’s frontier models—illustrates the new conundrum.

1. Anthropic vs. the Pentagon: It’s Not About Ethics, but Governance

In July 2025, Anthropic became the first frontier AI company authorized to deploy its models on classified US government networks—specifically, its Claude model. The agreement was notable because the Pentagon accepted Anthropic’s AI safety

guardrails, including restrictions on using its models for fully autonomous weapons and domestic mass surveillance. The relationship deteriorated in early 2026, however, when the Pentagon sought to renegotiate the contract and requested unrestricted use of Claude for any lawful military purpose. Anthropic refused to remove its safeguards, arguing that its redlines concerning these guardrails were fundamental to the company’s responsible AI commitments.¹

With Anthropic holding the line, on February 27, 2026, President Donald Trump directed all federal agencies to cease using Anthropic’s technology and gave most a six-month transition window. War Secretary Pete Hegseth simultaneously declared Anthropic a “supply chain risk,” a designation under the Federal Acquisition Supply Chain Security Act previously reserved for firms suspected of ties to foreign adversaries and never applied to an American company.² The designation required defense contractors to certify that they were not using Claude in connection with DoW contracts.

Anthropic challenged the decision in federal court, arguing that the designation was retaliatory and unlawful. The litigation has produced conflicting rulings: a federal judge in San Francisco issued an injunction in March to block enforcement of the designation, calling it “spectacular overreach”;³ the DC Circuit Court of Appeals, by contrast, denied Anthropic’s request for a stay, finding in April that “the equitable balance” favored the government’s interest in “how, and through whom, the Department of War secures vital AI technology during an active military conflict.”⁴ Anthropic remains excluded from new Department of War contracts while continuing to serve other federal agencies, and the underlying legal question—whether a contractual safety redline can be recast as a national security threat—remains unresolved.

Corporate vs. State Sovereignty

Public commentary has largely framed the Anthropic-Pentagon dispute as a disagreement over the ethical boundaries

of artificial intelligence. The episode immediately evoked comparisons with Google's 2018 decision to withdraw from Project Maven after employees protested over the military use of AI. However, framing the dispute primarily as an ethical disagreement misses the central issue. The real question is not where to draw the moral boundaries of AI, but who ultimately controls critical national security capabilities developed and operated by private companies.

The Pentagon's forceful response can be understood in this context. A military may view a supplier that reserves the right to suspend support during a crisis not as a contractor with strong ethical commitments, but as a potential operational vulnerability. By labeling Anthropic a potential supply chain risk, the Department of War signaled that reliability and continuity of service are strategic requirements. Its underlying concern was never primarily whether terms like *autonomous weapons*, *human in the loop*, or *mass surveillance* described desirable objectives, but who retains the authority to define, interpret, and enforce those terms over time. Allowing a private company to reserve the right to suspend services based on its own reading of contractual language would hand corporate executives a degree of operational authority traditionally reserved for elected civilian leaders and military commanders. The issue was not the substance of Anthropic's restrictions, but the precedent the arrangement would set.

The argument rests on three observations:

1. Concepts central to debates about AI ethics are far less precise than they appear. What constitutes an offensive use of AI? What qualifies as a weapon? When is a human meaningfully in the loop? These are not technical questions; they are legal, political, and strategic judgments that democratic societies traditionally assign to legislatures, courts, and military authorities, not to a vendor and its terms of service.

2. If individual firms impose distinct operational restrictions, governments face a fragmented landscape in which an operation might comply with one provider's policy while violating another's—a patchwork generating uncertainty precisely where clarity is most needed, at the moment of crisis.
3. As AI systems evolve from productivity tools into critical infrastructure, a provider's ability to deny or suspend access becomes a source of strategic leverage in and of itself. The question shifts from AI ethics to political authority: Should a private corporation possess a *de facto* veto over state action?

There is huge ambiguity over whether industry or government should ultimately set the rules for high-impact technologies. This does not imply that companies should have no role in shaping norms for AI use or that governments should enjoy unlimited discretion. There is an unresolved tension at the heart of the AI age: Advanced systems are developed by private actors but increasingly function as instruments of national power. The Pentagon's position, reduced to its essence, is that if AI becomes critical infrastructure for warfare, continuity and sovereign control matter as much as performance. Anthropic's position is that developers cannot disclaim responsibility for how customers use their systems. The unresolved question—where corporate responsibility ends and government authority begins—is arguably the most consequential governance issue in AI today. The Anthropic dispute, alongside the subsequent Mythos/Fable suspension, is unlikely to be the last case to expose it.

Where the Encyclical Crosses This Road

These reflections intersect directly with Pope Leo XIV's first encyclical, *Magnifica Humanitas: On Safeguarding the Human Person in the Time of Artificial Intelligence*, presented at the Vatican on May 15, 2026.⁵ The document is, among other things, a sustained meditation on the tension the Anthropic affair dramatizes: governance authority migrating from states to private, often transnational actors.

The pope is explicit on this point. “In the past,” he writes, “it was largely up to the State to guide and direct innovation. Today, however, the main drivers of development are private, often transnational, parties that are endowed with resources and the capacity to intervene that surpass those of many Governments.” The encyclical calls for AI to be “disarmed”—not in the sense of rejecting the technology, but, in the pope’s own formulation, by “discrediting the assumption that technical power automatically confers the right to govern.” Regulatory tools, he argues, are necessary but insufficient on their own; the prior question is who holds this power today and how they use it.

Although the encyclical does not single out particular companies or individuals, its critique clearly addresses concerns about the growing concentration of technological power in a small number of frontier AI firms and digital platforms. In the American context, it can be understood as an implicit challenge to the increasing influence exercised by major technology companies—and, more broadly, by the ecosystem of entrepreneurs and investors who advocate a limited role for public regulation while playing an expanding role in national security and defense innovation.

The encyclical implicitly supports approaches that seek to balance innovation with accountability and aligns with broader discussions surrounding global AI governance. It ultimately frames the AI debate through two competing metaphors: Babel, representing technological hubris, limitless power, and innovation detached from moral purpose, and Jerusalem, symbolizing cooperation, solidarity, and technology oriented toward human flourishing.

The pope’s formulation—power without an automatic mandate to govern—captures the theological aspect of the same problem that the Pentagon dispute raised in legal and strategic terms: A private actor’s technical capability does not, by itself, settle the question of who has the authority to decide how

that capability is used. The encyclical does not adjudicate the Anthropic case; notably, Anthropic co-founder Chris Olah was among the speakers at its Vatican presentation, where he argued that AI developers alone cannot determine the ethical boundaries of their own systems because competitive and financial incentives shape the technology—a concession from inside the industry that lends the pope’s broader argument unusual force.

For governments wrestling with civil-military authority over AI, and for an industry still negotiating where its own responsibility ends, *Magnifica Humanitas* supplies something the legal record cannot: vocabulary for the legitimacy question sitting beneath the contract dispute.

2. Mythos and Fable: The Perfect Example of US and Europe Decoupling

A second, distinct episode involving Anthropic compounded the dispute with the Pentagon. In April 2026, the company released its Mythos Preview model but limited access to it because it could identify and exploit software vulnerabilities. This capability set off alarm bells across Silicon Valley and Washington. On June 9, Anthropic released Mythos 5 and Fable 5, Fable being a safety-constrained public variant of the more powerful Mythos model, restricted in its cyber-offensive capabilities. Within one day, on June 12, the US government issued an export control directive ordering Anthropic to suspend all access to both models for any foreign national, anywhere in the world, including the company’s own non-citizen employees. Because Anthropic had no technical means of filtering access by nationality across its global customer base in real time, it disabled both models for all users, everywhere.

The company stated that the issue was limited, comparable to issues previously identified in other frontier models, including

OpenAI's GPT-5.5, and did not amount to a universal jailbreak. It also argued that the findings did not justify withdrawing access to a model already deployed to millions of users, and criticized the government's decision-making process as lacking transparency and adequate technical justification. However, the US government's decision to restrict access to Anthropic's frontier models reveals more than a temporary policy dispute: It exposes the emergence of competing visions of how frontier AI should be governed. As these models acquire strategic significance, governments and AI companies are advancing different—and sometimes conflicting—conceptions of authority, responsibility, and sovereignty.

The US Vision: Frontier Model AI as a Strategic and National Security Asset

Observers should understand these restrictions in the broader context of the strategic competition between the United States and China. They extend Washington's technology denial strategy beyond advanced semiconductors to frontier AI models, reflecting the growing perception that access to the most capable foundation models constitutes a strategic advantage comparable to access to cutting-edge chips or critical defense technologies.

A recent analysis by the Center for a New American Security argues that frontier AI will increasingly determine military effectiveness, cyber capabilities, intelligence analysis, scientific discovery, and long-term economic competitiveness.⁶ According to the analysis,

Chinese advanced artificial intelligence (AI) systems pose a serious and growing threat to US national security. At least seven Chinese developers now produce systems with formidable capabilities across coding, reasoning, multimodal recognition, and agentic tasks—systems that are released with open weights, offered via application programming interface (API) at prices designed to undercut

American competitors, and available for download by anyone in the world. The Chinese Communist Party (CCP), which collapses the boundaries between state, military, and private sector, treats these systems as instruments of political control, economic dominance, and great-power competition.⁷

From this perspective, preserving US leadership in frontier AI is not only an economic objective but also a national security imperative. The rapid progress of Chinese models—including DeepSeek and, more recently, Z.ai's GLM-5.2, which now approaches leading US frontier models on several coding and agentic benchmarks—has reinforced concerns in Washington that unrestricted access to American models could accelerate China's military and technological rise through techniques such as model distillation and knowledge transfer.

The concern extends beyond China's official research institutes. US policymakers increasingly fear that frontier models accessible through commercial APIs could facilitate technology transfer to Chinese companies, researchers, or state-affiliated organizations, whether they operate in China or abroad. Access restrictions are therefore intended to reduce the risk of strategic competitors replicating, adapting, or weaponizing US-developed advanced reasoning, coding, cyber, or scientific capabilities.

From this perspective, frontier AI models are no longer ordinary commercial software but strategic technologies whose diffusion governments need to manage, just as they do for advanced semiconductors or other dual-use capabilities. Therefore, in addition to mitigating immediate security risks, export controls preserve America's long-term technological, economic, and military advantages—even if they impose costs on allies and partners that rely on access to US frontier models. This is reflected in Executive Order 14409 of June 2, 2026, which requires American AI companies to provide the federal government with access to

their newest models before public deployment.⁸ Combined with recent restrictions on access to Anthropic’s models and a new form of government oversight placed on GPT-5.6, these measures suggest that frontier models are beginning to be treated less as commercial products than as dual-use technologies whose dissemination may be subject to national security considerations.

While some observers have compared this evolving posture to China’s system of state oversight over advanced technologies, the comparison should not be overstated: The legal foundations, institutional framework, and objectives remain fundamentally different. The rationale advanced by Washington is rooted in nonproliferation considerations and national security. If a frontier reasoning model is judged sufficiently capable—and potentially vulnerable to misuse or hostile distillation—the government increasingly considers limiting its dissemination to be exercising sovereign responsibility.

Europeans initially perceived the episode differently. Whereas Washington framed the restrictions as a matter of national security and technological nonproliferation, the immediate European reaction centered on a different concern: the sudden realization that access to a critical technology could be unilaterally suspended. Rather than becoming a debate on the dangers of unrestricted dissemination of frontier models, the discussion focused on the emergence of the long-feared “kill switch” and the vulnerability it exposed for Europe’s technological sovereignty.

The European Kill Switch Nightmare

A US government directive issued late on a Friday resulted in two of the world’s most capable AI models, Fable 5 and Mythos 5, becoming unavailable by the following morning, without any public explanation of the technical rationale behind the decision. Regardless of the merits of the government’s concerns, the episode demonstrated that political decisions in

Washington could almost instantaneously suspend access to frontier AI models.

For Europe, the implications were immediate. Access to frontier AI had long been assumed to be commercially stable and politically neutral. Instead, its suspension exposed a new strategic vulnerability: dependence on AI capabilities ultimately controlled by foreign governments and companies. The episode sent shockwaves across Europe as policymakers understood the vulnerability was not theoretical. The French Conseil de l’Intelligence Artificielle et du Numérique (Council for Artificial Intelligence and Digital Affairs), an official observatory for the French government, described the suspension as “the materialization of the kill switch, a risk that has been repeatedly highlighted in recent months.” According to the council, the episode confirms that “digital technologies are fundamentally an instrument of power” whose exercise is becoming increasingly direct and explicit.⁹ The scale of the European reaction also suggests that Washington underestimated how profoundly such a decision would reverberate among allies that depend on US frontier AI capabilities.

The implications extend well beyond Anthropic. Only days later, the Trump administration requested that OpenAI stagger the release of GPT-5.6, limiting initial access to a restricted group of approved partners while federal agencies conducted security evaluations. Together, these episodes demonstrate that frontier AI models are no longer ordinary digital services but strategic technologies whose availability may be subject to national security considerations.

European governments and organizations investing heavily in AI need to consider frontier AI models as critical digital infrastructure and adapt their risk assessments to this reality: A focus solely on cybersecurity or technical performance is not enough. They should also account for geopolitical dependence, portability, interoperability, vendor concentration, and the

possibility that access to essential AI capabilities could be restricted with little notice.

Competing Governance Paradigms for Frontier AI

All three actors—the United States, frontier AI companies, and Europe—now implicitly recognize that frontier AI models are strategic assets. Their disagreements no longer concern the importance of AI, but who should govern it and under which principles.

For the United States, frontier AI has become a strategic capability comparable to advanced semiconductors or other dual-use technologies. Because these models underpin military capability, scientific leadership, economic competitiveness, and intelligence, Washington increasingly considers them to be subject to sovereign authority, export controls, and national security oversight.

Private AI companies promote a different governance model. Firms such as Anthropic continue to frame frontier AI as a technology that should ultimately benefit humanity, while seeking to preserve a degree of autonomy over how their models are deployed and by whom. At the same time, they selectively grant privileged early access to trusted partners, develop strategic alliances with governments, and expand internationally through cloud partnerships, infrastructure investments, and enterprise deployment. Their ambition is therefore not only technological leadership but also the ability to shape the norms governing frontier AI.

Europe represents a third vision. Rather than controlling frontier models or relying on corporate self-governance, the European approach has focused on reducing technological dependencies, strengthening regulatory oversight, and fostering sovereign industrial capabilities. Yet recent events have exposed this strategy's limitations. As European governments and industries prepare to deploy AI at scale, they remain dependent on frontier models whose availability may

ultimately be determined by decisions taken in Washington or by a handful of private companies.

Only weeks after the US Commerce Department suspended access to Fable 5 and Mythos 5 for non-American users, it lifted the export restrictions, reportedly after Anthropic agreed to strengthened safeguards and a framework granting the US government access to future frontier models before their public release. This rapid reversal illustrates both the volatility of this new environment and the emergence of a different governance model. Whether this arrangement becomes the template for other frontier AI companies, including OpenAI with the anticipated release of GPT-5.6, remains to be seen.

The Anthropic episode therefore represents more than a dispute over export controls or AI safety. It marks the emergence of competing constitutional visions for governing frontier AI. The United States increasingly treats frontier models as strategic capabilities subject to sovereign authority. Meanwhile, frontier AI companies seek to preserve corporate discretion over their deployment, and Europe seeks to reconcile technological openness with strategic autonomy through regulation and industrial policy. The question is no longer simply who develops the most capable AI systems, but who has the legitimate authority to decide how they are governed, who may access them, or who may withdraw access to them and under what conditions.

Consequently, organizations and governments need to assess not only the performance of AI systems but also the geopolitical and legal conditions governing continued access to them.

3. The Myth of a Global AI: A Vision of Shared Frontier Models Collides with the Reality of Power Politics

For years, frontier AI companies promoted artificial intelligence as a universal project—building intelligence for the benefit of humanity, developing general-purpose models capable of

serving all people, and creating governance mechanisms that transcend national borders. Sam Altman frequently spoke of technology designed to “benefit humanity as a whole”;¹⁰ Anthropic championed an approach centered on safety, alignment, and global responsibility. Underlying both narratives was the assumption that sufficiently advanced artificial intelligence would constitute a global common good requiring unprecedented international cooperation.

The recent events have told a different story: The idea of neutral artificial intelligence is collapsing, along with the myth of a universal AI, developed in the general interest of humanity and accessible to all as a global public good. The dream of a shared intelligence for humanity is increasingly colliding with the logic of strategic competition.

If frontier AI models are considered strategic national assets, governments will seek to influence their development, secure access to computing resources, control data flows, and shape corporate technological choices. The companies themselves will negotiate directly with states, particularly their militaries and intelligence agencies. AI is ceasing to be a universal project and is becoming an instrument of power.

Anthropic has provided a particularly revealing example, and 2026 has supplied the clearest illustration yet. The company continues to advocate a vision grounded in safety and the public interest. Yet it found itself simultaneously blacklisted by the Pentagon as a supply chain risk and suing its own government in two federal courts, and three months later was forced into a global shutdown of its most advanced models, all while maintaining that its conduct was consistent with exactly the safety-first principles it had always claimed to represent. This is not a contradiction unique to one company. It illustrates a deeper structural condition: AI companies often claim to speak on behalf of humanity while necessarily operating within, and increasingly at the mercy of, specific national, legal, and geopolitical environments.

This contrast between the ambition and reality has four implications:

1. It marks the end of the illusion of purely global AI governance. International institutions will remain important, but the most consequential decisions are being made in Washington, Beijing, and somehow, in the future, in Brussels, London, and Paris, in the boardrooms of a handful of technology companies—and, as the export-control directive showed, sometimes in a single afternoon.
2. It heralds a new form of technological sovereignty. In the twentieth century, states sought control over energy resources, strategic industries, and defense production. In the twenty-first, they increasingly seek control over computational capacity, foundation models, and digital infrastructure—and, as the Mythos/Fable case shows, the ability to revoke access to all three without warning.
3. It challenges the assumption that foundation models will naturally converge toward universal values. Large-scale AI systems inevitably reflect cultural, political, and legal choices. The question is no longer simply how to align AI with human values, but with whose values and under whose authority—a question posed in explicitly civilizational terms by Pope Leo XIV’s encyclical, which asked not only what AI can do, but who holds the power to direct it and toward what end.
4. It poses a major challenge for Europe. If foundation models become instruments of power comparable to energy infrastructure or military capabilities, European dependence on American or Chinese AI ecosystems will not merely be an economic issue but will directly affect Europe’s ability to preserve strategic autonomy, normative preferences, and political sovereignty.

A year ago, the central debate revolved around the prospect of artificial general intelligence and the transformative potential of agentic AI. Today, the more consequential development is

political rather than technological: the emergence of a world in which states, frontier AI companies, and competing systems of values seek to shape rival AI ecosystems. This shift may prove to be one of the defining geopolitical transformations of the twenty-first century.

Palantir CEO Alex Karp, in a recent interview with the French journal *Le Grand Continent*,¹¹ argues that Europe risks confusing technological sovereignty with what he calls “techno-politicization”—making procurement and technological decisions for political rather than operational reasons. Excluding foreign providers may create the illusion of sovereignty while reducing military effectiveness. Yet the suspension of access to Anthropic’s frontier models illustrates the opposite side of the equation, demonstrating that dependencies on foreign-controlled AI systems can themselves become strategic vulnerabilities. If Europe has sometimes politicized the AI debate, recent events suggest that it was not entirely misguided in treating access to frontier AI as a geopolitical concern rather than as a mere commercial issue.

The anticipated initial public offerings of OpenAI and Anthropic reinforce this transition. They signal that frontier AI is entering a new phase in which the decisive challenge is no longer simply building the most capable model, but financing the enormous computing infrastructure required to train, deploy, and continuously improve it. The hundreds of billions of dollars expected to flow into frontier AI will not merely fund research; they will determine which companies can build the hyperscale data centers, secure access to advanced semiconductors, attract developers, and establish the platforms on which governments, businesses, and societies increasingly depend. The contest is therefore becoming one over infrastructure as much as intelligence.

At the same time, the ambition of making advanced AI an open and universally accessible technology is coming under increasing pressure. Rising development costs, growing national security restrictions, and intensifying geopolitical competition are challenging the AI-for-everyone promise that featured prominently at international AI summits in Paris and New Delhi.¹² Access to frontier models is unlikely to remain indefinitely subsidized, as publicly traded companies face increasing pressure to generate returns on unprecedented capital investments.

Together, these developments point toward a broader transformation: from AI as a widely available digital service to frontier AI as scarce strategic infrastructure controlled by a small number of companies operating under the authority—or increasingly the influence—of a handful of governments. In this emerging landscape, access to frontier AI may become a geopolitical privilege rather than a commercial service. The defining competition of the coming decade may therefore be less about who invents artificial intelligence than about who owns, finances, governs, and controls the infrastructure on which the world’s most powerful AI systems depend.

Anthropic, OpenAI, SpaceXAI, and Google DeepMind are no longer simply technology firms. They have become strategic institutions whose decisions affect national security, industrial competitiveness, democratic resilience, and even military planning. The governance of these companies, their relationships with governments, the distribution of authority within them, and the conditions under which states can influence or constrain them have therefore become geopolitical issues in their own right.

Endnotes

- 1 David Jeans et al., “Anthropic Digs in Heels in Dispute with Pentagon, Source Says,” Reuters, February 25, 2026, <https://www.reuters.com/world/anthropic-digs-heels-dispute-with-pentagon-source-says-2026-02-24/>.
- 2 Matt O’Brien and Konstantin Toropin, “Trump Orders Federal Agencies to Stop Using Anthropic Technology in Dispute over AI Safety,” *DefenseNews*, February 28, 2026, <https://www.defensenews.com/news/pentagon-congress/2026/02/27/trump-orders-federal-agencies-to-stop-using-anthropic-technology-in-dispute-over-ai-safety/>.
- 3 Michael Kunzelman, “Appeals Court Judges Appear to Be Divided over Pentagon’s Legal Dispute with AI Company Anthropic,” Associated Press, May 19, 2026, <https://apnews.com/article/ai-anthropic-trump-security-risk-a8cfd07b4d975ddfc5be7e016ed3ddce>.
- 4 *Anthropic PBC v. United States Department of War and Peter B. Hegseth*, No. 26-1049, US Court of Appeals for the District of Columbia Circuit, April 8, 2026, p. 7, https://storage.courtlistener.com/recap/gov.uscourts.cadc.42923/gov.uscourts.cadc.42923.01208838678.0_1.pdf.
- 5 Pope Leo XIV, *Magnifica Humanitas: On Safeguarding the Human Person in the Time of Artificial Intelligence* (Holy See, May 15, 2026), <https://www.vatican.va/content/leo-xiv/en/encyclicals/documents/20260515-magnifica-humanitas.html>.
- 6 Daniel Remler *Red Lines: Understanding the National Security Risks of China’s Advanced AI* (Center for a New American Security, June 2026), https://s3.us-east-1.amazonaws.com/files.cnas.org/documents/Red-Lines_TECH_Final.pdf.
- 7 Remler, *Red Lines*, 1.
- 8 “Promoting Advanced Artificial Intelligence Innovation and Security: Executive Order 14409,” White House, June 2, 2026, <https://www.whitehouse.gov/presidential-actions/2026/06/promoting-advanced-artificial-intelligence-innovation-and-security/>.
- 9 “Dependance Day ? Fable 5, Mythos 5 : l’Europe face à son point de bascule” [Dependence Day? Fable 5, Mythos 5: Europe facing its tipping point], Council for Artificial Intelligence and Digital Affairs, June 15, 2026, <https://www.conseil-ia-numerique.fr/nos-travaux/dependance-day-fable-5-mythos-5-leurope-face-son-point-de-bascule>.
- 10 “Introducing OpenAI,” OpenAI, December 11, 2015, <https://openai.com/index/introducing-openai/>.
- 11 “Palantir et l’illusion de la souveraineté” [Palantir and the illusion of sovereignty], *Le Grand Continent*, July 3, 2026, <https://legrandcontinent.eu/fr/2026/07/03/palantir-et-lillusion-de-la-souverainete-texte-commentaire>.
- 12 “AI for People – Paris,” Observer Research Foundation, January 16, 2026, <https://www.orfonline.org/event/ai-for-people-paris>; “New Delhi Summit Calls for AI in the Public Interest,” Agence Française de Développement, February 23, 2026, <https://www.afd.fr/en/news/new-delhi-summit-calls-ai-public-interest>.



About the Author



Tsiporah Fried is a visiting senior fellow at Hudson Institute, focused on transatlantic relations, European defense and military strategy, and defense and tech innovation.

Ms. Fried is a defense and strategy expert with over 15 years of experience in military strategy, net assessment, and future warfare analysis. Before joining Hudson, she served as senior advisor to the vice chairman of the French Joint Chiefs of Staff, where she led work on defense innovation, emerging technologies, and strategy. She has contributed to major assessments of global military capabilities and strategic trends, while coordinating forward-looking strategic studies and managing a broad network of experts across academia, think tanks, and agencies.

© 2026 Hudson Institute, Inc. All rights reserved.

About Hudson Institute

Hudson Institute is a research organization promoting American leadership for a secure, free, and prosperous future.

Founded in 1961 by strategist Herman Kahn, Hudson Institute challenges conventional thinking and helps manage strategic transitions to the future through interdisciplinary studies in defense, international relations, economics, energy, technology, culture, and law.

Hudson seeks to guide policymakers and global leaders in government and business through a robust program of publications, conferences, policy briefings, and recommendations.

For information about supporters of Hudson, please see pages 73–75 in our *Annual Report 2024*, available at <https://www.hudson.org/annual-report-2024>.

Visit www.hudson.org for more information.

Hudson Institute
1201 Pennsylvania Avenue, NW
Fourth Floor
Washington, DC 20004

+1.202.974.2400
info@hudson.org
www.hudson.org