



TEN PREDICTIONS REVISTED

AI Decrypted 2026: Mid-Year Assessment



Stay ahead with us.



Introduction

In January, we wrote that the central question was no longer whether models would improve, but who would control the infrastructure, energy, data, talent, and regulatory pathways that determine where artificial intelligence can be deployed at scale. Now, well into the second half of 2026, the central question has evolved into who can reliably assemble, finance, govern, and secure access to the necessary stack—and set the terms on which others use it.

With that in mind, we are reviewing the ten emerging trends and challenges we identified at the beginning of the year and how they will continue to define the future of AI heading into 2027. In short, we were largely correct in our assessments. Let's review why.

Prediction scorecard

#	PREDICTION	STATUS	KEY REVISION
1	AI bubble fears overblown	CORRECT	Agents and inference validated; If there is bubble, it is in financial arrangements funding some capex, impacting bond market
2	U.S. infrastructure buildout hits physical limits	CORRECT	Bottlenecks are real; Gulf capacity is additive, not yet a substitute. Moratoriums way up, becomes major political issue, slowing buildout
3	Federal AI initiatives stretch the bureaucracy	CORRECT	Mobilization exceeded expectations; conversion to outcomes remains unresolved.
4	Regional hubs and AI sovereignty take off	MOSTLY CORRECT	Multipolar deployment is real; frontier capability remains concentrated.
5	Global governance fully fractures	CORRECT	No progress at the international level after India AI Summit, U.S.-China Dialogue holds hope, but faces major challenges.
6	Open-versus-closed model clash escalates	CORRECT	Chinese open weights lead diffusion; the U.S. response remains chaotic, and administration considering regulation, punishment for distillation
7	China's alternative AI stack remains incomplete	CORRECT	Domestic capability improved quickly, but still heavy reliance on Nvidia for training, including via access outside China
8	The Brussels Effect hits an AI wall	MOSTLY CORRECT	Regulatory influence persists while the EU (European Union) pivots toward industrial policy.
9	U.S. federal-state regulation becomes a quagmire	CORRECT	Preemption efforts added a new layer of legal and compliance uncertainty.
10	Advanced AI diffusion raises incident risk	CORRECT	Warning indicators and policy responses rose; Fable/Mythos is a game changer.

AI Bubble Fears Overblown: Agents and Inference Deepen Compute Demand



CORRECT

January thesis: Agents, multimodal systems, world models, and embodied AI would move from demonstrations into practical use, creating a durable increase in training and inference demand.

Assessment

2026 has become the year AI shifts from generating content to consuming compute in pursuit of goals. Agents are the clearest near-term driver of that transition. The decisive constraint is increasingly not whether models can perform useful work but whether organizations can afford, verify, authorize, secure, and govern that work at scale – and justify the token costs.

Our January analysis was correct. Agents would move into real work, coding capability would accelerate sharply, inference demand would rise, consumer devices would become agentic, physical AI would expand, and security and governance would become first-order constraints.

In 2026, an important change has been the shift from short chatbot exchanges to delegated, long-horizon work. McKinsey's August [survey](#) found that 40% of respondents at organizations with more than \$1 billion in annual revenue were scaling AI agents in at least one function, up from 27% a year earlier. Adoption is strongest in software engineering, IT, and knowledge-management workflows, where outputs can be tested, reviewed, and reversed.

The compute mechanism is also clearer than it was in January. An agent does not simply answer once: It plans, reads context, calls tools, inspects results, retries failures, and may run for minutes or hours. OpenAI reports that agentic coding now accounts for a majority of combined ChatGPT and Codex output tokens among its enterprise customers. The International Energy Agency (IEA) likewise identifies agents and other complex inference workloads as an important reason aggregate electricity demand can keep rising even as energy use per task falls rapidly.

In addition, multi-agent systems have advanced substantially. Anthropic introduced agent teams in Claude Code in February, allowing multiple agents to work simultaneously on a software task. Google DeepMind published Co-Scientist in May as a multi-agent system in which specialized agents generate, criticize, rank, and refine scientific hypotheses. Google DeepMind has also launched a dedicated safety research program for a potential future environment containing millions of interacting agents.

Finally, the IEA reported in April that electricity consumption by AI-focused data centers rose approximately 50% in 2025. It expects AI-focused data-center power demand to triple by 2030 despite extraordinarily rapid efficiency improvements at the task level. Moreover, the IEA explicitly identifies reasoning, agentic tasks, and video generation as substantially more energy-intensive than simple text generation—potentially consuming hundreds or thousands of times more energy per query.

ADOPTION IS ACCELERATING

McKinsey's August survey found that

40%	vs.	27%
of respondents at organizations with more than \$1 billion in annual revenue were scaling AI agents in at least one function in 2026		in 2025

WHAT TO WATCH THROUGH 2027

The next phase of the compute boom will be driven by millions of agents consuming inference in pursuit of goals. This will include agents operating on consumer smartphones, though regulatory issues around deployment on these platforms remain unresolved. Coding will remain the leading edge because its outputs are unusually easy, to test. Bubble risk will lie in financial arrangements that funds capital expenditures, thereby further impacting the bond market.

U.S. AI Infrastructure Buildout Runs into Physical and Political Limits



CORRECT

January thesis: Power, equipment, and permitting limits would slow U.S. projects and push major model developers toward Gulf superclusters as a near-term relief valve.

Assessment

Our January analysis was correct that power, equipment, and domestic politics would emerge as initial binding constraints on U.S. AI infrastructure and that foreign compute would begin to introduce strategic dependencies.

The IEA reports tightening supply chains for transformers, gas turbines, advanced chips, and other datacenter components, while planning and approval systems are struggling with the volume of proposed projects. In June, the Federal Energy Regulatory Commission (FERC) ordered all six regional grid operators under its jurisdiction to justify or reform the tariffs governing data centers and other high-load usage. The reforms focus on interconnection studies, cost allocation, co-location, behind-the-meter generation, and flexible-load arrangements.

“Physical bottlenecks have not caused U.S. model developers to abandon domestic expansion. Instead, constrained firms are responding by vertically integrating power, infrastructure, financing, and community arrangements.”

What has not happened is a retreat from domestic construction. OpenAI reported in April that it had surpassed 10 gigawatts of planned capacity for U.S. Stargate and that GPT-5.5 had been trained at its operational campus in Abilene, Texas. We would note that planned capacity should not be confused with energized capacity. Even so, they show that constrained firms are responding by vertically integrating power, infrastructure, financing, and community arrangements rather than simply moving core workloads overseas.

As we predicted, the United Arab Emirates (UAE) has moved closer to the center of the U.S.-aligned stack when the Commerce Department placed it in Country Group A:5 designation of the Export Administration Regulations (EAR). This grants approval for license-free transfers of advanced computing items to the UAE government and selected entities but would be subject to safeguards and investment commitments. It remains unclear that Gulf campuses are absorbing a material share of U.S. frontier-training demand. The Gulf is becoming an additional node in the U.S. stack, but it has not yet been shown to be the relief valve for capacity that the United States cannot build.

In July, New York imposed the first statewide moratorium on hyperscale data center permits by executive order, and New Jersey now requires facilities drawing 50 megawatts or more to carry their full infrastructure cost. As November's midterm elections near, the intersection of AI infrastructure, energy policy, and local community impact is approaching a critical juncture. While well-capitalized firms are positioned to navigate higher costs and complex energy arrangements, the scale of their operations makes them focal points in debates over affordability, environmental impact, and grid reliability. How policymakers balance these competing priorities will play a key role in shaping the next phase of AI infrastructure development in the U.S.

WHAT TO WATCH THROUGH 2027

Physical bottlenecks have not caused U.S. model developers to abandon domestic expansion. But regional diversification also carries new risks. During the conflict with Iran this year, AWS reported sustained disruption to its Bahrain Region and advised affected customers to migrate workloads to other regions. Gulf capacity can hedge U.S. grid and permitting constraints, but it cannot be treated as a risk-free substitute. Finally, domestic politics surrounding data center moratorium efforts have the potential to exacerbate polarization of the electorate to a degree that creates significant investor risk in AI.

Trump Administration AI Initiatives: Mobilization Outpaces Conversion



CORRECT

January thesis: The AI Action Plan, Genesis Mission, Pax Silica, U.S. Tech Force, and related initiatives would move into execution but collide quickly with limited personnel, fragmented authority, and interagency resistance.

Assessment

In eight months, the Trump administration has achieved significant programmatic execution. Genesis expanded from a Department of Energy (DOE) initiative into a governmentwide AI-for-science program involving more than 15 agencies. The White House announced more than \$5 billion in federal commitments; DOE selected 278 initial projects and reported more than \$800 million in partner support. Pax Silica expanded to 24 signatories, while the U.S. Tech Force added industry partners and specialized recruitment pathways.

The institutional model has been more federated than centralized. The White House sets direction, while agencies with existing authorities oversee different parts of the stack: DOE leads AI for science and energy infrastructure; General Services Administration (GSA) and Office of Management and Budget (OMB) handle procurement; Office of Personnel Management (OPM) runs talent programs; Department of State, Department of Commerce, and finance agencies coordinate exports and supply chains; National Institute of Standards and Technology (NIST) and Center for AI Standards and Innovation (CAISI) handle measurement and standards; and Federal Energy Regulatory Commission (FERC) addresses grid rules. This structure reduces the need for a large new AI bureaucracy inside the White House.

That does not mean that our original concern was misplaced. The execution test has moved downstream. A selected project is not a scientific result; a funding commitment is not money spent; a federal-land lease is not an operating data center; and a temporary technical corps is not durable state capacity. Genesis will still confront data rights, classification, privacy, cybersecurity, and verification constraints as teams seek to combine sensitive datasets and move from model output to reproducible science.

In addition, on frontier AI models specifically, the advent of Mythos in early 2026, and the loss-of containment episode in July headlined by the OpenAI agentic platform hack of Hugging Face and proposals for a new AI self-regulatory organization (SRO) galvanized efforts to finally establish a quasi-government oversight function for the most advanced models, an idea much discussed at the end of the previous administration. Major questions about its structure remain, including whether congressional authorization will be needed.

In sum, despite being formally launched with much spectacle early in 2026, most of these initiatives have failed to move the needle at the market level. With no new sources of materials yet operational, it remains unclear whether the U.S. has a comprehensive long-term strategy to meaningfully reduce dependence on Chinese firms across a wide range of critical minerals, particularly heavy rare-earth elements and the finished projects incorporating them.

WHAT TO WATCH THROUGH 2027

The right metric for the remainder of 2026 and 2027 is conversion: the share of announced programs that become funded projects, working systems, energized capacity, validated discoveries, and permanent institutional capability. Another key development to watch is a potential executive order mandating establishment of an AI SRO, a process likely to take 12 to 18 months in the best-case scenario.

Regional AI Development Takes Off, but the Frontier Remains Concentrated

MOSTLY
CORRECT

January thesis: The Middle East, Japan, India, and Europe would become major AI hubs as governments built domestic compute, local models, and sovereign deployment pathways.

Assessment

Our January analysis was largely correct: 2026 would accelerate the early shift toward a more geographically distributed AI ecosystem, and AI sovereignty would become one of the defining concepts of global technology policy.

Governments are increasingly building domestic compute pools, sovereign clouds, local-language models, procurement frameworks, and sector-specific deployment environments, even as advanced accelerators, fabrication capacity, hyperscale clouds and most leading models remain concentrated in a small number of countries and companies.

A Geographically Distributed Ecosystem—Different Stages of Development

JAPAN	EUROPE	INDIA	THE GULF
Japan's SoftBank is building a domestic GPU cloud, while the country's Digital Agency is testing Japanese foundation models on a government cloud.	Europe has launched a call for as many as seven AI Gigafactories alongside its existing AI Factory network and sovereign cloud procurement.	India has onboarded more than 38,000 GPUs into a shared national compute pool and is backing domestic models trained on Indian data and languages.	The Gulf, meanwhile, remains the most ambitious sovereign capital and energy story, but large frontier-training capacity is still mostly nascent.

In 2026, sovereign AI emerged as a distinct market category. Companies such as Reflection AI, Mistral, Cohere/Aleph Alpha, G42/Core42, HUMAIN, Cerebras, Nscale and Armada are offering governments and national champions ways to run advanced AI with greater control over model weights, data, compute, deployment and operations. The market ranges from model-centric sovereignty—Reflection and Mistral—to sovereign compute providers such as Cerebras and Nscale, and full-stack national platforms such as G42 and HUMAIN. This reflects a broader effort to avoid permanent dependence on U.S. or Chinese hyperscalers and closed-model providers.

WHAT TO WATCH THROUGH 2027

Regional hubs are becoming sovereign interfaces to a still-concentrated global AI stack. The competitive advantage will go to providers that can move workloads across jurisdictions while preserving data controls, security, and application behavior.

Global Governance Fragments into Layers Rather than Fully Fracturing



CORRECT

January thesis: U.S.-China competition would divide AI governance into rival blocs, weaken the Bletchley summit process, and make global interoperability increasingly difficult.

Assessment

Eight months in, fragmentation has accelerated but has not fully fractured. Global AI governance is fragmenting into layers, not breaking apart. Instead, it is becoming a layered regime in which universal dialogue, technical cooperation, national regulation, and strategic technology alliances coexist. The U.S. and China increasingly lead competing ecosystems around the physical and economic AI stack, but middle powers continue to participate selectively in both, while the UN, India, Switzerland and technical institutions preserve a thin layer of interoperability.

The U.S. and China are building increasingly distinct technology and governance networks. Pax Silica expanded to 24 signatories and links AI to critical minerals, energy, advanced manufacturing, semiconductors, and infrastructure. China formally launched the World Artificial Intelligence Cooperation Organization (WAICO) in Shanghai in July with 29 founding countries and a mandate covering cooperation, governance, standards, and capacity building.

At the same time, broad multilateral coordination did not collapse. The New Delhi AI Impact Summit declaration was endorsed by 92 countries and international organizations. The United Nations held the inaugural Global Dialogue on AI Governance in Geneva, supported by an independent scientific panel. The summit sequence that began at Bletchley Park is continuing, even as its agenda expands from frontier safety toward development, infrastructure, labor, access, and sovereignty.

The emerging system is layered. Universal forums preserve high-level principles. Technical networks cooperate on measurement, evaluations, and security science. National and regional regulators impose different market rules. Strategic blocs compete intensely over chips, compute, infrastructure, capital, and supply chains. The closer governance gets to control of scarce physical assets, the more geopolitical it becomes; the closer it gets to technical measurement, the more cooperation survives.

Middle powers are not choosing one complete package. India, Gulf states, Southeast Asian governments, and others can align with Washington on hardware, use Chinese models or development programs, adopt European regulatory concepts in selected sectors, and participate in UN processes simultaneously. The result is competitive multi-alignment, not two sealed camps.

The debate in Washington has intensified over how to move quickly on these issues, in part to reassure investors and markets and demonstrate U.S. leadership on frontier model AI governance. At the same time, developments around Mythos, cybersecurity and loss of containment have pushed the U.S. and China to consider closer cooperation on AI governance. But major hurdles remain to reaching an agreement and developing a framework that could gain international acceptance.

WHAT TO WATCH THROUGH 2027

All of this sets the stage for the U.S.-China Dialogue on AI happening this week in Washington, D.C. While there are significant opportunities for collaboration, both sides lack clarity around how their own governments are postured for frontier AI governance. Industry calls for greater government capacity to regulate frontier AI will pick up, putting more pressure on both the U.S. and China, and efforts to gain international agreement in the runup to next year's AI safety summit in Switzerland.

Open-versus-Closed Model Competition Becomes a Dual Ecosystem



CORRECT

January thesis: U.S. labs would release more open-weight models, but Chinese labs would lead global diffusion and force policymakers to treat model weights as strategic assets.

Assessment

Our January thesis was directionally correct. The most capable U.S. frontier systems remain predominantly closed, while open-weight models have become practical tools for local deployment, fine-tuning, sovereign clouds, research, and lower-cost applications.

The Trump administration's AI Action Plan explicitly encourages open-source and open-weight AI models. There are now reports and policy discussions concerning a potential executive order targeting Chinese open-weight or open-source models used in the U.S. Such an order would need to fit several converging developments: rapidly expanding U.S. company and researcher adoption of frontier models from Alibaba, DeepSeek, Z.ai (formerly Zhipu - Wisdom Spectrum), Moonshot, and other Chinese firms, including for production applications; concerns about data leakage, censorship and remote dependencies; allegations that Chinese developers distilled U.S. models; and the Fable/Mythos realization that advanced models increasingly resemble strategic cyber capabilities rather than ordinary software.

Chinese authorities are simultaneously considering restrictions on foreign access to their own most advanced models, partly in response to U.S. treatment of Mythos, reinforcing the likelihood of a more explicitly reciprocal "model controls" competition.

Chinese labs, meanwhile, have strengthened their position in that open ecosystem. Hugging Face's August review found that Alibaba's Qwen models have become one of the community's most important base-model families, with an unusually large derivative ecosystem. DeepSeek V4, Z.ai's GLM, Moonshot's Kimi, and other Chinese model families have expanded the range of capable systems available under permissive or relatively deployable terms. Their advantage is not only benchmark performance; it is coverage across model sizes, languages, and hardware budgets.

The U.S. response is increasingly dual-track. The AI Action Plan explicitly supports open-source and open-weight development, and U.S. companies such as Nvidia and startups such as Thinking Machines continue to invest in open models. At the same time, Washington is treating unauthorized distillation, foreign model dependencies, and access to advanced weights as national security questions. An April White House memorandum alleged industrial-scale distillation campaigns by foreign entities, principally in China, while emphasizing that legitimate distillation and open-weight development remain important to the ecosystem.

This is creating divisions within the Trump administration over how to regulate Chinese AI models while setting up a new process of pre-release evaluation of closed frontier AI models. The administration is also forming a new organization to regulate the release of frontier models, with significant lingering questions on the future shape of regulation.

WHAT TO WATCH THROUGH 2027

Expect more debate over provenance, distillation, government procurement, and whether some advanced weights should face access controls. Going forward, movement on issuing guidance covering government and contractor use; certain sectoral use; and having bodies such as the Center for AI Standards and Innovation, which evaluates Chinese models, issue further reporting around potential risks the models pose. Critical to the entire effort will be the timeline for development of a U.S. AI SRO, which could be the ultimate arbiter to U.S. policy on open weight models.

China's Alternative AI Stack Advances, but Is Not Fully Independent

**CORRECT**

January thesis: Chinese firms would combine available Western-origin hardware with rapid investment in domestic accelerators, software, and model ecosystems, but the alternative stack would remain incomplete.

Assessment

Our January analysis was correct that power, equipment, and domestic politics would emerge as initial binding constraints on U.S. AI infrastructure and that foreign compute would begin to introduce strategic dependencies.

Chinese open-weight model families are globally competitive in many practical uses, and domestic accelerator systems are becoming more credible. In July, Huawei introduced the Atlas 950 SuperPoD as a system-level approach to training and high-concurrency inference, emphasizing large clusters, proprietary interconnects, and an increasingly open domestic software stack. As production of the Ascend 950 process ramps towards the end of this year, there will be more real-world deployments of domestic hardware systems designed for training advanced models.

“China is best understood as building a viable parallel stack rather than a fully self-sufficient one.”

System-level engineering matters because the contest is not decided by one chip specification. China can offset weaker individual accelerators through scale, networking, memory architecture, software optimization, and model efficiency. Here, Huawei's Tau (τ) Scaling Law appears to provide a real advantage and a way to get to more advanced performance levels on par with 3 nanometer western systems. The diffusion of Qwen and other Chinese models also reduces the software ecosystem disadvantage by attracting global developers and derivative work.

The gap has not disappeared. NIST's May evaluation assessed DeepSeek V4 Pro as the strongest Chinese model it had tested but estimated that its aggregate capabilities lagged the U.S. frontier by about eight months. Since then, frontier models from Z.ai, Alibaba, and Moonshot have narrowed the gap further. Advanced lithography, leading-edge fabrication yields, high-bandwidth memory, power efficiency, and the cost of migrating software from CUDA-centered workflows remain structural constraints. Huawei's own performance claims should therefore be read as evidence of rapid progress, not independent proof of parity. In June, reports suggested that a domestic immersion lithography capability may be farther along than some had anticipated, though major hurdles appear to remain, including system downtime and long-term support in production environments.

China is best understood as building a viable parallel stack rather than a fully self-sufficient one. Its likely near-term architecture is hybrid: Domestic hardware and software take a growing share of strategically important workloads, while firms continue to exploit any lawful access to foreign components, tools, and research.

WHAT TO WATCH THROUGH 2027

The most important indicator is not whether a single Chinese accelerator matches single U.S. chip. It is whether Chinese systems can deliver competitive training and inference economics at cluster scale, with reliable supply, mature tooling, and acceptable energy costs.

The Brussels Effect Meets Industrial Reality

MOSTLY
CORRECT

January thesis: Europe would discover the limits of rule-setting without frontier compute and would pivot from regulation toward public investment, sovereign cloud, and industrial policy.

Assessment

Six months after we published AI Decrypted, Brussels' focus has shifted from regulating U.S. technology platforms to fostering European Union (EU) infrastructure and services.

The EU's promise to buy U.S. advanced chips to train AI models is one of the few that will likely hold — simply because the bloc has no viable alternative. Some 96% of AI chips sold in the EU already come from U.S.-based vendors, including Nvidia, while the bloc has no AI chip designers or manufacturers of its own.

This summer, the EU executive will launch a process to build seven massive AI computing hubs, the three largest of which will need at least 40,000 AI chips, while the four smaller ones will need at least 25,000. Some European startups are trying to build alternatives, but those efforts require ample time and funding. That means Europe will need to buy U.S. AI chips for the foreseeable future.

The EU remains a consequential market whose rules affect product design, documentation, transparency, and deployment. The European Commission (EC) began enforcing the AI Act and new transparency requirements on August 2. At the same time, the July AI Omnibus extended several high-risk-system deadlines, simplified obligations for smaller companies, and expanded testing mechanisms. Europe is trying to preserve the credibility of the AI Act while reducing the risk that compliance schedules outrun technical standards and industrial capacity. The proposed AI Liability Directive was withdrawn, so the original report's suggestion of a new strict AI-liability regime should be removed.

Industrial policy has become much more prominent. Europe now has a network of 19 AI Factories, has launched a tender for as many as seven larger AI Gigafactories, and has proposed the Cloud and AI Development Act (CADA), to expand datacenter capacity, accelerate deployment, and create a common framework for cloud and AI sovereignty.

Europe nevertheless remains dependent on foreign-designed accelerators, global semiconductor manufacturing, and non-European cloud technology. Public compute and procurement can reduce dependency at the deployment layer, but they cannot quickly create a domestic frontier hardware industry.

Mistral has emerged as Europe's clearest vertically integrated AI champion, expanding beyond model development into the infrastructure layer needed to train and serve those models. Alongside its frontier and specialized models, Mistral now offers Mistral Compute, providing European GPU capacity, training and inference services, regional APIs, and private or on-prem deployments, with a stated goal of securing up to 1 gigawatt of sovereign European AI capacity by 2030. The result is a European alternative not only to OpenAI and Anthropic at the model layer but increasingly to the U.S. hyperscalers at the infrastructure and deployment layer as well.

WHAT TO WATCH THROUGH 2027

Europe's influence will increasingly depend on whether regulation is paired with usable compute, sovereign deployment options, procurement demand, and globally competitive firms.

U.S. Federal-State Regulatory Tensions Deepen

**CORRECT**

January thesis: The absence of a comprehensive federal AI statute would produce growing conflict among states, Congress, regulators, and the executive branch, increasing compliance uncertainty.

Assessment

Our predictions were accurate. States continue to legislate automated decisions, transparency, discrimination, child safety, deepfakes, employment, health, and frontier-model obligations, while Congress has not enacted a durable national framework. The administration has responded through executive order (EO) rather than waiting for legislation.

Federal AI legislation in 2026 remains fragmented and is well short of a comprehensive national law. The Trump administration issued a national legislative framework in March centered on innovation, workforce and federal preemption of burdensome state AI laws, but Congress remains divided over preemption and the appropriate level of safety regulation; meanwhile, narrower bipartisan bills covering model documentation, security, research, deepfakes and federal AI use have advanced through committees.

A December 2025 EO (14365) created an AI Litigation Task Force to challenge selected state laws, directed Commerce to evaluate state measures, and asked the Federal Communications Commission (FCC) and Federal Trade Commission (FTC) to explore federal reporting and consumer protection theories that could preempt conflicting state requirements. In July, the FTC published a proposed policy statement explaining how state mandates that alter model outputs could conflict with the FTC Act's prohibition on deceptive practices.

These actions may reduce fragmentation if courts uphold them, but in the near term they add another layer of uncertainty. Companies should treat federal preemption as a contingency to monitor rather than a plan to build against — and note that even enacted preemption would run three years, not permanently, which makes it a deferral of the patchwork rather than a resolution of it. Organizations should also track state-effective dates while assessing whether federal agencies will challenge a law, whether funding conditions are enforceable, and whether courts will accept broad preemption theories. The result is not simply fifty different state codes; it is overlapping state, federal, sectoral, and judicial processes with uncertain boundaries.

EO 14365 of December 11, 2025, published as 90 FR 58499, asserted federal authority to preempt state artificial intelligence laws and established an AI Litigation Task Force to challenge them. On March 20, the White House followed with a national policy framework asking Congress to codify that approach.

The preemption question is being litigated and negotiated rather than legislated, and the calendar makes legislating it before 2027 improbable. The only comprehensive vehicle has not been introduced, its sponsors face opposition from their own chamber's AI leadership, and the floor between now and November belongs to appropriations.

WHAT TO WATCH THROUGH 2027

The defining question in federal AI policy is not what to regulate but who gets to regulate it. Congress has declined to enact preemption more than once, having rejected it in reconciliation and again in the defense authorization. Until Congress acts or courts clarify preemption, state rules will continue to shape deployment even when national policy opposes them.

Advanced AI Diffusion Raises Risk

**CORRECT**

January thesis: As advanced models and agents diffuse into sensitive domains, visible AI-mediated harm would increase and governments would begin constructing frontier-model security and incident-response frameworks.

Assessment

We were spot on here. The emergence of advanced cyber capabilities, first with Mythos in early 2026, and then a series of other models, including OpenAI's Astra, brought the challenges of releasing advanced models with national security related capabilities to the fore. The June Fable/Mythos testings and release was also highly impactful. According to Anthropic, the U.S. government directed the company to suspend foreign-national access to the models on national-security grounds three days after release. Anthropic disabled them globally because it could not implement the restriction immediately; the controls were lifted later that month and access was restored on July 1. Whatever one's view of the merits, the episode exposed the absence of a stable, transparent process for emergency model-access decisions.

The administration also issued Executive Order 14409 on June 2. It directs federal agencies to strengthen cyber defense, expand access to advanced defensive tools, develop classified capability benchmarks, and create a voluntary process through which developers can provide the government limited pre-release access to covered frontier models. A new U.S.-China channel on AI governance also opened after the May leaders' summit, creating at least the possibility of narrow risk-reduction discussions amid strategic competition.

The events of the past several months, including the emergence of advanced models, such as Anthropic's Mythos, will substantiate cybersecurity related capabilities, the uncertainty surrounding the release of a guardrailed version of Mythos in June, and then, in another huge inflection point, the spate loss of containment incidents, headlined by the OpenAI agentic platform hack of Hugging Face have all accelerated calls for the US government to establish an orderly process for model release, as well as galvanized efforts to establish new government bodies to oversee frontier AI model releases. While cybersecurity is now front and center because of Mythos, other issues such as biosecurity remain important, along with loss of control.

Finally, in May, June, and July, a series of blogs and letters from industry researchers highlighted growing concerns within the leading AI labs that recursive self-improvement (RSI) could lead to rapid increases in model capability and require governments to develop mechanisms for pausing or pacing development of advanced AI platforms.

WHAT TO WATCH THROUGH 2027

The next policy trigger is more likely to be a concrete demonstration of consequential autonomous capability, a serious jailbreak, or an operational incident than another benchmark gain. Governments will need clearer due process, common severity scales, and channels that work at incident speed. There is now a race for developing government capacity to gate models and if the need arises to coordinate—including globally and in particular U.S. and Chinese government agreement—on slowing of model development as the benefits of RSI kick in.

The potential for a malicious actor to leverage a frontier model and inflict major virtual or physical damage on critical systems is rising rapidly in the absence of clear government capacity to oversee and control model development and multi-agentic deployments. This could push the AI safety debate more strongly toward the national security and increase consideration of more government nationalization of advanced AI labs.

About Us

DGA Group is a global advisory firm that helps clients protect – and grow – what they have built in today's complex business environment. We understand the challenges and opportunities of an increasingly regulated and interconnected world. Leveraging the experience and expertise of **Albright Stonebridge Group**, a leader in global strategy and commercial diplomacy, and a deep bench of communications, public affairs, and government relations consultants, we help clients navigate and shape global policy, reputational and financial issues. To learn more, visit dgagroup.com.



Please contact **Paul Triolo**, Partner and Technology Policy Lead, at Paul.Triolo@dgagroup.com, and **Alexis Serfaty**, Managing Director, at Alexis.Serfaty@dgagroup.com, with any questions or to arrange a follow up conversation.

