

Toward a Federal Framework

Lessons from State and International Frontier AI Regulation

By Laura Caroli and Aalok Mehta

Executive Summary

In 2023—the year after ChatGPT ushered in the generative artificial intelligence (AI) era—safety was top of mind for many policymakers and AI developers around the world. OpenAI openly advocated for the licensing of powerful frontier models, the European Union scrambled to modify its AI Act to address generative models, and world governments assembled for a landmark AI Safety Summit. Today, the political winds have shifted dramatically. The U.S. federal government, concerned about overregulation and its chilling effects on AI development, has actively opposed attempts by U.S. states to regulate frontier AI, and international convenings have shifted away from considerations of safety and governance to a focus on sovereignty, investment, and adoption, despite recent developments in model capabilities that are driving renewed attention to risks.

Yet, notwithstanding the political shifts, there are signs of an emerging consensus on regulating frontier AI models, focused on implementing and publishing safety frameworks. For example, U.S. state AI bills—such as California’s **Transparency in Frontier Artificial Intelligence Act** (S.B. 53), New York’s **Responsible AI Safety and Education** (RAISE) Act, **Illinois S.B. 315**, **Michigan H.B. 4668**, and **Massachusetts S. 2630**—share core priorities with international approaches, and their core elements are being taken as inspiration for a possible **federal bill**. This consensus reflects a complex coevolution between governments and companies. While governments must weigh the trade-offs between protecting citizens, fostering innovation, and competing economically and technologically with other countries, companies must bear the costs of regulation and acknowledge their role in building legitimacy and trust. However, this consensus is imperfect. While the broad contours may be coming into focus, countries around the world, as well as U.S. states, disagree on key definitions and questions of scope—details that may seem small but have historically proved to be significant obstacles to legislative action.

This report compares several U.S. state bills to international approaches to regulating and shaping frontier AI development, including the EU AI Act, and voluntary industry governance frameworks

for frontier AI models. The analysis, which extends through June 2026, reveals that U.S. states are serving as laboratories of democracy by trialing a variety of approaches to AI legislation. The most pro-regulatory versions of these experiments have failed to gain traction, whereas more modest approaches that mirror industrial practice have become law. These attempts could provide a blueprint for the U.S. government's stated goal of putting into place a federal legislative framework that can preempt a state patchwork of laws and reduce the compliance burden for U.S. companies.

FINDING 1: CONVERGENCE ON HOW TO GOVERN A MODEL

The comparison of U.S. state bills and international approaches reveals four measures for frontier model developers that appear repeatedly in U.S. and international legislation, forming an emerging global baseline for responsible frontier AI oversight:

- **Safety framework:** implementing and regularly updating a safety and security protocol (SSP) or comparable framework outlining risks, mitigation strategies, and security measures for a set of severe harms
- **Transparency:** publicly disclosing the SSP, often in redacted form
- **Risk assessment:** identifying the AI model's potentially dangerous capabilities
- **Risk mitigation:** creating measures to prevent or reduce catastrophic outcomes

Besides these four pillars of frontier AI governance, nearly all the examined frameworks also require security measures, such as protection of model weights and systems against unauthorized access. Whistleblower protections are also common across most frameworks, with the exception of RAISE and the Seoul Frontier AI Safety Commitments. In addition, many frameworks include incident reporting requirements, though reporting timelines vary from 72 hours (RAISE) to tiered approaches based on severity (California S.B. 53 and the EU Code of Practice for General-Purpose AI Models).

FINDING 2: DIVERGENCE ON WHAT TO GOVERN

While governance approaches converge across the different frameworks, their scope differs in four aspects:

1. **Value-chain coverage:** The type of AI operators that bear obligations varies. RAISE and California S.B. 53 impose obligations only on model developers, whereas the now-defunct California S.B. 1047 placed obligations on developers, compute providers, and even auditors.
2. **Size of covered models:** U.S. bills often use a high training threshold to define covered entities, as measured in FLOPs (floating-point operations per second), inspired by U.S. President Joe Biden's [Executive Order 14110](#), and compute cost or revenue thresholds. EU thresholds, by contrast, are lower and allow direct designation of smaller models posing a systemic risk.
3. **Risk categories:** All U.S. bills cover chemical, biological, radiological, and nuclear (CBRN) risks, serious crimes with minimal human control, and cyberattacks; most add loss of control over AI. The European Union also adds societal risks such as large-scale manipulation or rights violations.
4. **Catastrophic risk definitions:** All current U.S. bills identify catastrophic risk using high casualty or economic thresholds (more than 50-100 deaths and more than \$1 billion in damages). Conversely, the European Union takes a broader approach that includes infrastructure disruption, rights violations, and environmental harm.

Overall, the EU regime covers more models, more risks, and a wider set of harms than any current U.S. state law.

WHAT THESE FINDINGS MEAN FOR U.S. AI POLICY

States have introduced hundreds of AI bills in the past few years, covering a wide range of topics. Bills aimed at regulating frontier models have taken varying approaches. The heaviest regulatory approaches have failed to gain traction, whereas more moderate approaches have become law. In this sense, U.S. states are serving as important test beds for AI policy regulation. This exploration can—and should—inform federal approaches instead of being seen as an impediment.

Allowing state frameworks to evolve and leveraging their innovations at the federal level could, in theory, accelerate the path to national AI legislation, which would directly address concerns over state fragmentation while reducing national security and economic risks from AI misuse and enhancing trust in the U.S. tech stack globally. This approach is consistent with the goals of key U.S. policy constituencies and would advance goals articulated in [America's AI Action Plan](#), which remains the guiding strategy for U.S. AI policy.

Introduction

From the outset of President Donald Trump's second term, his administration has been consistent and deliberate in redirecting U.S. artificial intelligence (AI) policy toward deregulation. Departing sharply from the approach of President Joe Biden's administration—epitomized by [Executive Order 14110](#) and its focus on governing AI risks and ensuring safety—the Trump administration has channeled its energy into ensuring that the United States prevails in the global AI race, particularly against China. This strategic pivot centers on dismantling perceived obstacles to AI innovation, including regulation deemed burdensome. This direction is set out clearly across several key policy documents: [America's AI Action Plan](#), released in July 2025; a [December 2025 executive order](#) on “Ensuring a National Policy Framework for Artificial Intelligence”; and the White House's [legislative recommendations](#) on the national AI policy framework, released in March 2026. Yet despite this federal approach, U.S. states have continued advancing and enacting AI-related legislation, including in areas of particular concern for the administration, such as model development.

[America's AI Action Plan](#) opens by calling for the removal of “red tape and onerous regulation.” Two of the plan's recommendations are particularly relevant in this regard, signaling the administration's ongoing focus on preempting state laws it sees as harmful to AI innovation:

- **Limit AI funding to states with onerous AI regulations.** The plan recommends the Office of Management and Budget work with federal agencies with “AI-related discretionary funding programs to ensure . . . they consider a state's AI regulatory climate when making funding decisions and limit funding if the state's AI regulatory regimes may hinder the effectiveness of that funding or award.”
- **Conduct a Federal Communications Commission (FCC) review of state AI regulations.** The plan recommends the FCC “evaluate whether state AI regulations interfere with the agency's ability to carry out its obligations and authorities.”

The [December 2025 executive order](#) confirms that these issues remain a priority for the administration. Among other provisions, the order creates an AI litigation task force to challenge AI state laws and outlines potential restrictions on federal funding to states with “onerous AI laws.” The

subsequent **legislative recommendations** on the national AI policy framework, a lean document outlining the elements of a potential federal AI law in line with the goals of the administration, identify seven pillars of AI, including “establishing a federal policy framework, preempting cumbersome state AI laws.” The goal of the framework is a light-touch national legislative framework that would include preemption of existing and future state AI bills.

Despite the Trump administration’s opposition to state AI laws, several elements of the policy vision outlined in the AI Action Plan align with recent work from U.S. state legislators. For example, the plan prioritizes policy actions to “invest in AI interpretability, control, and robustness breakthroughs,” “build an AI evaluations ecosystem,” and “ensure that the U.S. government is at the forefront of evaluating national security risks in frontier models”—all goals seen in various U.S. state bills. These are not novel policy actions; rather, they uphold past U.S. policies and are broadly consistent with emerging international approaches to frontier AI governance. However, few concrete details have emerged about how the administration plans to implement these provisions.

The administration’s deregulatory posture is still firm, but has its limits—and those limits become visible precisely when a frontier model crosses a capability threshold that existing voluntary frameworks were not designed to anticipate.

The administration’s concern for frontier AI safety has increased significantly following the **restricted release** of Anthropic’s Mythos model in April 2026. Mythos’s **reported** ability to autonomously identify thousands of software vulnerabilities across major operating systems prompted the White House to briefly explore, for the first time under this administration, a form of **predeployment government review** for the most capable AI systems. National Economic Council Director Kevin Hassett **stated** the administration is “studying possibly an executive order” to ensure that future AI models with significant security implications “go through a process so that . . . they’re released in the wild after they’ve been proven safe.” This draft executive order, ultimately released in June, established only a **voluntary framework** enabling participating companies to provide the federal government with predeployment access to their models. In parallel, the Center for AI Standards and Innovation (CAISI) entered **predeployment evaluation agreements** with Google DeepMind, Microsoft, and xAI, in addition to existing agreements with Anthropic and OpenAI, to assess frontier AI capabilities and national security implications before public release. In addition, on June 12, the federal government issued an export control **directive** blocking access to both Mythos and Fable to non-U.S. citizens.

These developments, including the hesitation over the mandatory or voluntary nature of the predeployment framework, confirm that the administration’s deregulatory posture is still firm, but has its limits—and those limits become visible precisely when a frontier model crosses a capability threshold that existing voluntary frameworks were not designed to anticipate.

State Legislation on Frontier AI

In the absence of a federal AI framework, several U.S. states have introduced legislation to ensure the safety and security of the most powerful AI models. Many frontier AI bills display a similar pattern of governance measures recalling the principles and measures seen in other international frameworks on frontier AI. These state bills include California S.B. 53, New York’s Responsible AI Safety and Education (RAISE) Act, Illinois S.B. 315, Michigan H.B. 4668, and Massachusetts S. 2630.

This report analyzes the main elements of these various U.S. state proposals and compares them to existing frameworks addressing similar issues, such as the EU AI Act, enacted in June 2024, and the Seoul Frontier AI Safety Commitments, signed by leading AI companies in May 2024. In addition, this report assesses the **Frontier Model Transparency Framework** proposed by Anthropic to evaluate how well government and multistakeholder approaches align with emerging industry best practices. Together, these bills align in key ways on governance of the most powerful AI models, including provisions that overlap with some elements of the EU AI Act and its Code of Practice for General-Purpose AI Models. This analysis provides insights into how U.S. state legislators’ thinking has evolved, how they might learn from past developments, how they can align with the international consensus on frontier AI governance, and what this implies for U.S. AI policy.

While Congress was discussing—and **ultimately dropping**—a 10-year moratorium on state AI legislation contained in the One Big Beautiful Bill Act, the New York Senate **approved the RAISE Act**. The act, passed in June 2025 and signed by Governor Kathy Hochul that December, regulates powerful AI models to ensure their safety and security and prevent malicious use. Meanwhile, in **September 2025**, California Governor Gavin Newsom signed the **Transparency in Frontier Artificial Intelligence Act**, referred to here as S.B. 53. The bill’s main proponent, Senator Scott Wiener, was also the author of the controversial **S.B. 1047**, which garnered significant industry pushback and **was vetoed** by Governor Newsom in September 2024. Both S.B. 53 and RAISE take a much lighter approach to frontier AI governance than S.B. 1047. In particular, they require developers to disclose basic due diligence measures to prevent catastrophic harm stemming from their frontier AI models.

Other state frontier AI bills that feature a similar structure include the recently approved **Illinois S.B. 315** and bills in earlier stages of discussion, such as **Michigan H.B. 4668** and **Massachusetts S. 2630**. These bills share many similarities with RAISE and S.B. 53. All introduce similar light-touch governance requirements, though each has its own nuances.

To reveal patterns across the policy landscape, this report compares nine frameworks on frontier AI governance: six U.S. state bills, including California S.B. 1047, New York’s RAISE Act, Illinois S.B. 315, Michigan H.B. 4668, California S.B. 53, and Massachusetts S. 2630; the Safety and Security Chapter of the EU AI Act Code of Practice; the Seoul Frontier AI Safety Commitments; and Anthropic’s **Frontier Model Transparency Framework** proposal. Anthropic’s proposal was selected over other existing company frameworks, such as those by **Google** or **OpenAI**, because it was published around the same time as the core discussions on key U.S. state bills were taking place and was explicitly intended to inform those debates. Indeed, contrary to other industry frameworks that leading AI companies have implemented internally and shared with the public, Anthropic **positioned its framework** as a recommendation for legislators “that could be applied at the federal, state, or international level, and which applies only to the largest AI systems and developers while establishing clear disclosure requirements for safety practices.”

A comparison of these nine approaches, shown in Table 1, underscores the central role of safety and security in frontier AI policy, as well as other important areas of convergence and divergence.

Table 1: Comparison of Frontier AI Model Developer Obligations Across Nine Frameworks

Obligation	California S.B. 1047	California S.B. 53	New York RAISE Act	Illinois S.B. 315	Michigan H.B. 4668	Massachusetts S. 2630	EU AI Act / Code of Practice	Seoul Frontier AI Safety Commitments	Anthropic Frontier Model Transparency Framework
Security measures	Yes	Yes*	Yes*	Yes*	Yes*	Yes*	Yes	Yes	No
Full shutdown enabled	Yes	No**	No	No	No	No**	No	No	No
Safety and security protocol (SSP)	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Internal governance measures	Yes	Yes*	Yes*	Yes*	No	Yes*	Yes	Yes	Yes
Document retention (at least five years)	Yes	Yes	Yes	Yes	Yes	Yes	Yes	No	Yes
Transparency	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Risk assessment	Yes	Yes*	Yes*	Yes*	Yes	Yes*	Yes	Yes	Yes*
Record testing procedures	Yes	No**	No**	No	Yes	No**	Yes	No	No
Risk mitigation	Yes	Yes*	Yes*	Yes*	Yes*	Yes*	Yes	Yes	Yes*
Refrain from deployment	Yes	Yes***	No**	No	Optional	No**	Yes	Yes	No
SSP review	Yes	Yes	Yes	Yes	Yes*	Yes	Yes	Implicit	Implicit
Third-party evaluation	Yes	Yes***	Yes***	Yes	Yes	Yes***	In most cases yes	Optional	No
Incident reporting	Yes	Yes	Yes	Yes	Yes*	Yes	Yes	No	No
Whistleblower protections	Yes	Yes	No**	Yes	Yes	Yes	Yes	No	Yes
Risk threshold definitions	No	Yes*	Yes*	Yes*	Yes*	Yes*	Yes	Yes	No
Possibility of excluding models from SSP (limited risk)	No	No	No	No	Yes	No	Concept of similarly safe model	No	Implicit

Note: * The obligation needs to be included in the safety and security protocol. ** The obligation was present in a previous version, then removed. *** The developer decides (and discloses what it did).

Source: CSIS.

Trends in Frontier AI Policy

SAFETY AND SECURITY ARE CONSISTENT ACROSS FRAMEWORKS

The concept of AI safety—once a prominent feature of national and international debates around AI policy, from Biden’s 2023 **executive order** on “Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence” to the 2023 UK **AI Safety Summit**—has faded to the background of U.S. and international policy language. A prominent example is the renaming of the third international AI summit, held in Paris in February 2025, to the **AI Action Summit**—a trend that continued when India held the 2026 iteration as the AI Impact Summit. The rebranding of the UK AI Safety Institute to the **AI Security Institute** and the renaming of the U.S. AI Safety Institute to the **Center for AI Standards and Innovation** demonstrate a similar trend.

The content of RAISE, California S.B. 53, and other U.S. state proposals demonstrates that AI safety and security remain at the core of proposals to regulate AI model development.

Nevertheless, the content of RAISE, California S.B. 53, and other U.S. state proposals demonstrates that AI safety (i.e., managing the risks stemming from powerful AI models) and security (i.e., ensuring robust cyber and physical security measures against malicious use) remain at the core of proposals to regulate AI model development. Indeed, there appears to be a rough consensus around broad elements of AI governance. Further, the **AI Action Plan**, despite its clear stance on deregulation, also maintains the need for safe and secure AI models, and it does so from the very introduction: “We must prevent our advanced technologies from being misused or stolen by malicious actors as well as monitor for emerging and unforeseen risks from AI. Doing so will require constant vigilance.”

Four obligations are consistent across all policies:

1. **Implementing a safety and security protocol (SSP) or framework.** A pivotal measure, the SSP is a document outlining the overall governance and compliance strategy of the companies involved. It typically includes other elements such as risk assessments, mitigation practices, and security protections.
2. **Performing a risk assessment.** A risk assessment requires a developer to evaluate the risks a specific model might pose when deployed. As shown later in this article, these risks vary across frameworks, from CBRN-related risks to broader societal risks.
3. **Implementing risk mitigation measures.** Once risks are identified, developers must determine whether and when to apply mitigation measures in order to prevent harm.
4. **Creating transparency by publishing the SSP.** Developers must publish their SSP, though they may do so in a redacted or summarized version to protect sensitive commercial information. At a minimum, the SSP must include mitigation measures to prevent critical incidents and to rapidly intervene when risk reaches unacceptable levels.

SECURITY MEASURES ARE WIDELY ACCEPTED

Security measures are another element present in all frameworks except Anthropic's. These measures consist of physical or cybersecurity policies to prevent malicious actors or unauthorized personnel from accessing the model and to protect the model's weights. Weights are numerical parameters that scale the signals passing between neurons in a neural network, determining how inputs are transformed step by step into the network's final output. Given the importance of protecting such a strategic asset for AI developers, it is surprising that these measures are missing from Anthropic's framework.

Such measures also seem to be in line with the spirit of the AI Action Plan, given its focus on national security against adversarial attacks. The word "security" appears 55 times in the plan, and the pillar on AI infrastructure has chapters on bolstering cybersecurity for critical infrastructure and promoting "secure-by-design AI technologies and applications." While these recommendations are aimed at securing government applications, the government also collaborates with most leading AI companies, meaning sound cybersecurity and physical security practices become necessary in every leading company's governance strategy.

Industry-Led Frameworks Use a Lighter Touch

The frameworks proposing the lightest-touch measures are the [Seoul Frontier AI Safety Commitments](#) and Anthropic's framework. The former, signed by 20 leading AI companies, likely reflects an iterative negotiation process with AI firms in the weeks before the summit. Endorsing a light-touch framework does not necessarily mean that companies want to minimize safety and security measures; rather, it indicates that industry prefers flexible frameworks and the possibility of retaining the independence to make decisions based on companies' internal characteristics, products, and business models.

THE KILL SWITCH PROVISION HAS LOST TRACTION

Obligations related to the potential full shutdown of models or derivatives (i.e., copies or plainly modified versions of the models) if they exceed certain risk thresholds do not appear to have significant stakeholder support. The kill switch, originally proposed in California S.B. 1047, was one of the bill's most [controversial provisions](#) and extended well beyond provisions considered during negotiation of the EU AI Act. None of the newer bills in California, Illinois, Michigan, or New York contain this measure. This highlights the complex dynamics of state lawmaking, which reflects and is shaped by significant stakeholder input, including feedback from AI companies.

NO CLEAR AGREEMENT ON THE ROLE OF THIRD-PARTY EVALUATIONS

When Governor Newsom vetoed S.B. 1047, he commissioned a [report on frontier AI policy](#) to provide policymakers with evidence-based reflections to shape future legislative efforts. The [California Report on Frontier AI Policy](#) was released in June 2025, and it highlights transparency as a key enabler to advancing research around frontier AI safety. For example, it proposes the use of third-party model evaluations, which involve independent experts testing model performance, safety characteristics, and vulnerabilities.

The state bills analyzed here take inconsistent approaches to the issue of mandatory third-party model evaluations for frontier model developers. For example, in RAISE, the obligation appeared in the [initial version](#) of the bill but faced fierce opposition and was removed in the signed text. In S.B. 53,

the obligation to retain a third-party evaluator **did not appear** at first; it was **added** in July 2025 and then removed again in the enrolled text, indicating a lack of clear agreement on how to approach the issue. The final text of California S.B. 53 makes the use of third-party auditing optional, only asking developers to disclose whether they used such a measure. The same solution was found in the final version of RAISE and in Massachusetts S. 2630. Recently approved Illinois S.B. 315 includes it as a hard obligation, as does Michigan H.B. 4668. But because Michigan H.B. 4668 is not close to final approval, this obligation may meet the same fate as it did in California S.B. 53.

The third-party provision is also controversial in Europe. The text of the **EU AI Act** does not mandate third-party evaluations; it considers them a valuable option. In fact, the act introduces only four obligations for frontier AI model developers (or “general-purpose AI model providers”): (1) model evaluation (internal or external), (2) risk assessment and mitigation, (3) incident reporting, and (4) cybersecurity. These obligations are in addition to the two baseline obligations that apply to all general-purpose AI model providers (including smaller ones), which involve technical documentation and copyright policy.

However, the four EU AI Act obligations for frontier AI model developers are broadly worded, with no details about specific compliance measures. Instead, legislators have opted for a flexible coregulatory framework due to the nascent nature of generative AI technology: the **Code of Practice for General-Purpose AI Models**.

The Code of Practice is a voluntary document detailing the measures that signatory companies commit to uphold. While the **final version** mandates third-party evaluations, it gives providers the option to avoid using a third-party evaluator when the model has a safe equivalent (implying it is similar to an existing one that is considered safe enough and therefore does not require a specific evaluation) or when the provider cannot find adequate experts. This provision, even in its final, somewhat reduced scope, was the subject of **heated debate**. The industry stated it went beyond the provisions of the EU AI Act, while prominent AI researchers such as Yoshua Bengio—a “godfather of AI” who serves as chair of the code’s **working group** on technical risk mitigation—insist that companies should not “**grade their own homework**.”

The appearance of third-party evaluations in various iterations of U.S. state AI bills could mean that some state legislators are, in principle, comfortable with exploring more aggressive regulatory and transparency regimes than EU policymakers, despite the latter’s reputation for broad and strict regulation. However, RAISE and S.B. 53 demonstrate the fragility of including such measures as legislation advances amid industry opposition. Third-party evaluation obligations were eventually kept in the Illinois bill, but they might still be removed or softened in the Michigan bill, which is in earlier stages of negotiation.

WHISTLEBLOWER PROTECTIONS ARE WIDELY ACCEPTED COMPONENTS OF A HEALTHY RISK ENVIRONMENT

A second measure suggested in Newsom’s AI report is whistleblower protections, which allow employees of model developers to inform government authorities if they become aware of practices within the company that might harm the public or violate the law, without fear of retaliation by the employer. Among all the analyzed frameworks, only the Seoul commitments and RAISE fail to mention

these protections (they appeared in the **initial version** of RAISE’s text but were later removed).

Whistleblower protections are **commonly accepted** in frontier AI **safety research** as highly effective but low-cost components of a sound frontier AI governance framework, which encourages early warning and reporting of breaches and safety risks. Their removal from RAISE represents a missed opportunity to align with international calls and improve overall safety.

Whistleblower protections are commonly accepted in frontier AI safety research as highly effective but low-cost components of a sound frontier AI governance framework, which encourages early warning and reporting of breaches and safety risks.

Although the other U.S. state bills and the EU AI Act contain this measure, the details vary widely. Both California S.B. 1047 and S.B. 53 detail how to inform employees of the available protections, such as by posting and displaying notices in workplaces and providing written notices annually. In S.B. 1047, “employees” included contractors and subcontractors working on models derived from other, more powerful models. S.B. 1047 also proposed strict retention periods for reporting (at least seven years).

In this regard, the EU AI Act has been **criticized** for providing few details on whistleblower protections. The **Code of Practice** lists only basic measures as examples of a “healthy risk culture.” However, the act falls within the scope of the **EU whistleblowing directive**, which details several provisions related to whistleblowers. The examples mentioned in the Code of Practice, therefore, serve as reinforcements of the existing regime in Europe, which is already **quite strict** in protecting whistleblowers.

RETENTION, RISK THRESHOLDS, REPORTING, AND REVIEW APPEAR FREQUENTLY

Beyond the obligations described above, other technical measures, such as document retention, setting up risk thresholds internally, SSP updates, and incident reporting, appear in most of the legislation considered here but not necessarily in the Seoul commitments or the Anthropic framework. These common measures include the following:

- **Document retention:** Developers are generally required to keep compliance documentation for the life cycle of the model plus five years. The EU AI Act requires 10 years.
- **Defining risk thresholds:** Developers must disclose how they assessed the risks associated with the frontier model, breaking down those risks into risk tiers within which the risk is considered acceptable.
- **Regular review of the SSP:** Developers must regularly update and review the SSP (generally at least once a year) to keep it current with a model’s evolving capabilities.
- **Incident reporting:** Developers must report incidents to the authorities within a specific period, with a description of the incident and the measures taken. Here, timelines vary according to the specific bill.

INCIDENT REPORTING PERIODS VARY WIDELY

Defunct California S.B. 1047 required developers to report incidents within the first 72 hours, which is sometimes used as a reporting time frame for **data breaches** or cybersecurity incidents affecting **critical infrastructure**. RAISE took the same approach at first but added to that a 24-hour time frame in case of imminent threat of death or serious physical injury. In Michigan H.B. 4668, the developer may decide whether and how to report incidents. The EU AI Act Code of Practice, California S.B. 53, Illinois S.B. 315, and Massachusetts S. 2630 use a more graduated approach, tying the timing of the reporting to the severity of the incident. In all three U.S. bills, the typical timing is 15 days, but the period is shortened to 24 hours in case of imminent threat of death or serious physical injury, as in RAISE. The EU Code of Practice adds even more tiers:

- Two days for a serious disruption of critical infrastructure
- Five days for a serious cybersecurity breach
- Ten days for the death of a person
- Fifteen days for harm to a person's health, fundamental rights, or property, or to the environment

Similar Frameworks, Varying Scope

Despite the many similarities seen in basic AI governance measures, the frameworks analyzed here demonstrate very different understandings of frontier AI policy. There is wide variation in the actors bearing obligations, the definition of a frontier AI model, the risks developers must control for, and even the definition of a catastrophic risk. These differences have a substantial impact on the number of models the frameworks cover and the extent of the testing and mitigation measures required to prevent catastrophic risks from materializing.

The federated nature of state policymaking in the United States lends itself to both pragmatic ideas and wild experimentation.

The federated nature of state policymaking in the United States lends itself to both pragmatic ideas and wild experimentation. For example, California S.B. 1047 would likely never have been introduced in Europe, let alone find a viable pathway to enactment. S.B. 1047 went much further than the EU AI Act in regulating what the latter calls “general-purpose AI models,” so much so that even staunch European supporters of AI safety would have considered it too burdensome and overbearing at the time, experts told CSIS. Indeed, EU AI Act negotiators—particularly representatives of EU member state governments—were reluctant to impose obligations on frontier AI providers from the outset. Farther-reaching measures, such as mandatory shutdown authority or supply chain obligations extending to compute providers, were never entertained by legislators, even those who advocated for the most stringent regulatory approach, especially in the European Parliament.

OPERATORS

This dynamic is particularly evident in how S.B. 1047 tackled the AI value chain—that is, the various operators involved in AI's development and diffusion. The bill would have imposed obligations on model

developers, operators of computing clusters, and third-party auditors. Both model developers and the operators providing them access to infrastructure to train their models would have “know-your-customer” obligations requiring them to exercise control over the models and their derivatives.

By comparison, Illinois S.B. 315 and Michigan H.B. 4668 introduce obligations only for covered models and third-party evaluators, and the EU AI Act imposes obligations only on model providers, including light transparency obligations for fine-tuners (i.e., those who retrain a frontier model with specialized datasets to create further applications). The lightest-touch U.S. state bills are RAISE, California S.B. 53, and Massachusetts S. 2630, which impose obligations only on frontier model developers. The differences among the frameworks are shown in Table 2.

Table 2: Value Chain Operators Covered in California S.B. 1047, California S.B. 53, New York’s RAISE, Illinois S.B. 315, Michigan H.B. 4668, Massachusetts S. 2630, and the EU AI Act

Operator	CA S.B. 1047	CA S.B. 53	RAISE	IL S.B. 315	MI H.B. 4668	MA S. 2630	EU AI Act
Computing cluster operators	Yes	No	No	No	No	No	No
Large model developers	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Third-party auditors	Yes	No (removed)	No	Yes	Yes	No	No
Downstream and derivative model providers	Yes	No	No	No	No	No	Basic transparency obligations

Source: CSIS.

COVERED MODELS

The size of the covered model (and the amount of computational power, or “compute,” used to train it) plays a key role in defining the scope of these bills. In this regard, they take varying approaches. California S.B. 1047 covered models trained with over 10^{26} FLOPs costing at least \$100 million in compute and fine-tuned models of at least 10^{25} FLOPs costing at least \$10 million. By contrast, the EU AI Act imposes basic obligations (technical documentation and copyright policy) on all models attaining at least **10²³ FLOPs** and the most burdensome obligations (model evaluation, risk assessment and mitigation, reporting of incidents, and cybersecurity) on models reaching at least 10^{25} FLOPs or those explicitly designated by the European AI Office.

Currently, Epoch AI, which provides estimates of the most notable large AI models currently on the market, **tracks** three models over 10^{26} FLOPs (Grok 3, Grok 4, and GPT 4.5) and around fifteen between 10^{25} and 10^{26} FLOPs as of May 2026. Thus, in practice, the EU AI Act would consider all of these in the

highest tier, triggering the frontier AI obligations. S.B. 1047, by contrast, would have covered fewer models due to the additional criterion of the training cost in both of its tiers (baseline and fine-tuned models).

A previous version of RAISE used similar criteria for the first tier as S.B. 1047 (10^{26} FLOPs, costing at least \$100 million). Its second tier comprised distilled models instead of fine-tuned models with a compute cost threshold of at least \$5 million. The mention of distilled models was likely added because of the concerns sparked by DeepSeek and its release of DeepSeek-R1, a distilled model, in January 2025. But while some concerns around it may be warranted, referring explicitly to distillation might not solve the problem if other techniques emerge that lead to small but very powerful models that entail safety challenges similar to those of DeepSeek-R1. Under that version of RAISE, these small models might have escaped the bill's rules because they are not a product of distillation. In the end, the New York legislators likely understood these implications and removed any reference to distillation from the [final text](#), thus converging with the approach taken by California S.B. 53.

The other state bills examined in this article take a somewhat different approach. California S.B. 53 takes a much lighter approach, covering only models above 10^{26} FLOPs whose developer's annual gross revenue was at least \$500 million in the preceding calendar year. The final version of RAISE and Massachusetts S. 2630 both ended up adopting the same criterion. In practice, this would cover only the three first-tier models and their developers—[xAI](#), as developer of Grok 3 and Grok 4, and [OpenAI](#), as developer of GPT 4.5—and only because both project revenues in the order of billions in the upcoming years (S.B. 53 would apply beginning January 2027).

Michigan H.B. 4668 takes yet another route. For models to fall under the scope of either bill, developers (1) must have spent over \$5 million in compute for their development and (2) must have developed over \$100 million worth of models in total computing cost over the previous 12 months. In practice, similar to the previous version of RAISE, these criteria seem to be aimed at covering DeepSeek-R1, a cost-efficient mini model purported to have cost only [\\$6 million](#), though the overall training compute costs that led to it seem to have been [much higher](#). By including the criterion of \$100 million spent over the previous year, Illinois H.B. 3506 and Michigan H.B. 4668 could aim at considering the process that led to DeepSeek-R1's development, which entailed hundreds of experiments culminating in the model's final training run. However, such coverage of cheaper models could be scrapped from this bill, as was the case with RAISE.

ARE FIXED FLOP THRESHOLDS FUTURE PROOF?

Given the speed at which frontier AI is evolving, it is important to preserve the possibility of updating thresholds. For instance, models at the 10^{26} threshold might go for months without serious incidents, pushing the frontier out to models trained at 10^{27} FLOPs or above. Conversely, new developments in model efficiency might provide frontier capabilities to smaller models, which policymakers may want to then capture under their frameworks. In general, the frameworks tracked here include the possibility of updating their compute threshold as AI technology evolves.

The EU AI Act allows the European AI Office to directly designate a model below 10^{25} FLOPs as posing systemic risk (and therefore triggering the top-tier obligations). A similar but less flexible model was proposed and later abandoned in [California S.B. 53](#):

If the Attorney General determines that less well-resourced developers, or developers significantly behind the frontier of artificial intelligence, may create substantial catastrophic

risk, the Attorney General shall promptly submit a report to the Legislature . . . with a proposal for managing this source of catastrophic risk but shall not include those developers within the definition of “large developer” without authorization in subsequently enacted legislation.

Therefore, California’s attorney general would have had powers like those of the EU AI Office to designate a small model developer as covered by S.B. 53, but only after obtaining legal authorization. This potentially would have represented a much longer, albeit more democratic, procedure than the more direct one allowed by the EU AI Act and could have assuaged the concerns Newsom expressed in [vetoing S.B. 1047](#):

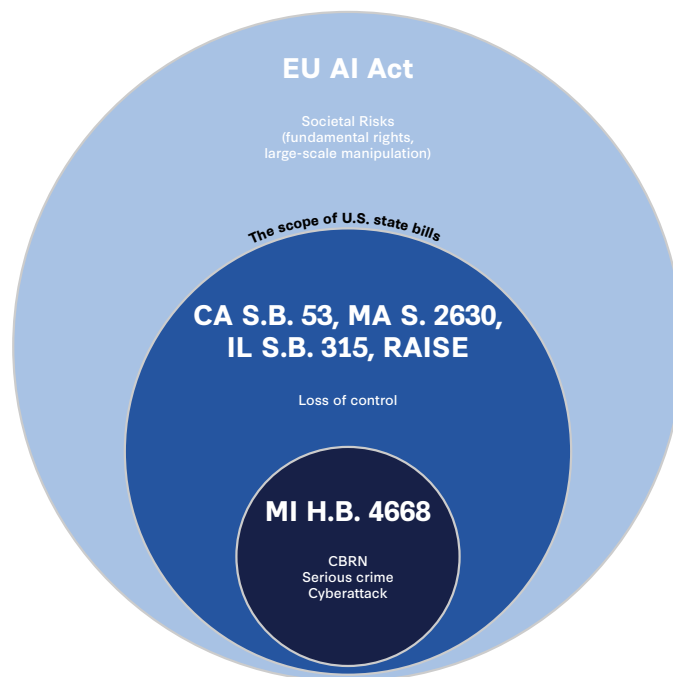
By focusing only on the most expensive and large-scale models, S.B. 1047 establishes a regulatory framework that could give the public a false sense of security about controlling this fast-moving technology. Smaller, specialized models may emerge as equally or even more dangerous than the models targeted by SB 1047.

In any case, the provision could provide a more flexible framework for Michigan H.B. 4668, as it would allow [designating](#) R1 as a covered model risk without incurring the risk of not being future proof. Perhaps the proponents of this bill could consider including it in their respective text and removing the lower threshold described above.

SCOPE OF RISK TAXONOMIES AND CATASTROPHIC IMPACTS

In developing frontier AI legislation, policymakers must decide which risks they aim to control, as this determines both how they protect their citizens and how burdensome compliance will be for companies. Here, more differences emerge than similarities, as illustrated in Figure 1.

Figure 1: Levels of Risk Considered Across Legislation



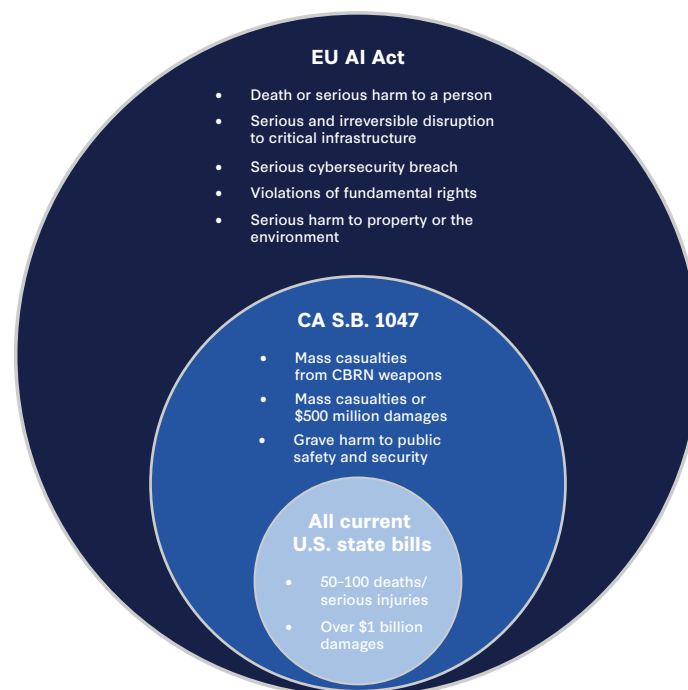
Source: CSIS.

Three risks form the “nucleus” of commonly accepted risks in the United States. All tracked U.S. bills consider at least the possibility of a model creating a CBRN weapon, as well as the possibility of a model committing a serious crime or a cyberattack with limited human control. More expansive are risks around the loss of control, added first by California S.B. 53 and later by Massachusetts S. 2630, Illinois S.B. 315, and RAISE in its final chapter amendment. But this is where the scope of U.S. state bills stops. The EU AI Code of Practice adds one further layer: societal risks, such as risks to fundamental rights and the risk of large-scale manipulation.

Figure 2 shows a similar comparison regarding the consideration of catastrophic impacts. All current U.S. frontier AI bills consider a narrow set of cases: incidents causing at least 100 deaths (reduced to 50 in the **enrolled text** of California S.B. 53 and, later, in RAISE, Illinois S.B. 315, and Massachusetts S. 2630) and more than \$1 billion in damages. California S.B. 1047 took a broader approach by including the following incidents:

- mass casualties from CBRN weapons;
- mass casualties or at least \$500 million in damages; and
- grave harm to public safety and security.

Figure 2: Critical Safety Incidents Covered in the Analyzed Frameworks



Source: CSIS.

U.S. legislators therefore seem to be moving away from the broad definition of catastrophic risk in California S.B. 1047 toward a narrower, more targeted definition of risk. The EU AI Act goes the furthest in defining a “serious incident,” which includes

- death or serious harm to a person;

- serious and irreversible disruption to critical infrastructure;
- a serious cybersecurity breach;
- violations of fundamental rights; and
- serious harm to property or the environment.

Thus, model developers must control for more risks if they are marketing their model in the European Union than if they are operating solely in the United States. They must also be prepared to report a broader range of incidents, though they have a longer time frame than under most U.S. bills.

CONVERGENCE ON THE HOW, DIVERGENCE ON THE WHAT

While the analyzed frameworks contain several notable similarities in governance processes, their scope (the models covered and the definition of a risk or a catastrophic incident) reveals substantial differences. The gap between the European Union and the United States is particularly evident. Indeed, the EU AI Act covers more models, more risks, and a broader set of incidents than most U.S. state bills.

Key Takeaways

This analysis reveals a strong degree of consensus on elements of basic frontier AI governance across U.S. states, the EU AI Act and its Code of Practice, the Seoul Frontier AI Safety Commitments, and Anthropic’s Frontier Model Transparency Framework. The measures with the broadest support across these frameworks include model developers establishing and publishing an SSP and conducting risk assessment and risk mitigation. Together, these measures constitute a foundational baseline of responsible frontier AI oversight.

COEVOLUTION EXPLAINS CONVERGENCE

Although EU policy may have influenced the other AI frameworks analyzed here, **the convergence likely does not stem from the European Union exporting its model** through a “**Brussels effect**,” to borrow professor Anu Bradford’s famous expression. If anything, the first internationally influential document about frontier AI governance came not from Europe but from the United States: President Biden’s **Executive Order 14110**. The order set out important principles that still underpin the field of frontier AI safety. For example, the compute threshold to determine the frontier of AI as specified in the order remains the reference point of most U.S. state bills and became a benchmark for EU AI Act negotiators, who were approaching the chapter on frontier AI at around the same time. The order was released around the time of the UK AI Safety Summit and its **Bletchley Declaration**, which also signified a commitment by global powers to pursue risk-based policies.

Instead, **all frameworks have evolved together** through parallel cross-pollinating processes involving similar stakeholders, from industry to internationally renowned academics, AI experts, and civil society organizations, working between the fall of 2023 and 2025. For example, Bengio was a **vocal supporter** of California S.B. 1047, and shortly after it was vetoed, he started his work as the chair of the **technical risk mitigation working group** of the EU Code of Practice. Similarly, Stanford researcher Rishi Bommasani is both vice chair of the EU Code of Practice transparency working group and one of the coauthors of *The California Report on Frontier AI Policy*. Key state lawmakers within the United States have also **coordinated legislative efforts**, and many of the same civil society organizations have been involved in both the EU Code of Practice platform and U.S. state bills (in some cases, they

were funded by the same philanthropic concerns). As the frontier AI industry is highly concentrated, only a handful of firms have the computing resources, technical expertise, and fundraising capability to continue innovation in model sizes and capabilities. Therefore, the companies voicing their concerns and providing input to regulators (including Meta, Google DeepMind, Anthropic, and OpenAI) are also largely the same across all frameworks.

In sum, this convergence of people and processes, not simply of policies, is happening simultaneously in multiple jurisdictions. Conspicuously absent from this process are voices from the Global South. However, a discussion of the reasons for and implications of this absence is beyond the scope of this report.

Major frontier AI developers already reflect this convergence on policy by implementing many of the same measures in their governance. For example, a [2025 report](#) by researchers at the Oxford Martin AI Governance Initiative at the University of Oxford analyzed several public industry frameworks and compared them to the EU AI Act Code of Practice measures for safety and security, finding a great deal of alignment. As Table 1 of this article shows, the code's provisions also map closely with requirements in U.S. state legislation and other international commitments. This reinforces the Oxford researchers' conclusion that there is meaningful global coherence in emerging frontier AI governance regimes.

WOULD THESE MEASURES BE ENOUGH?

While this convergence suggests a path forward on international harmonization of frontier AI laws, the most accepted measures represent the minimum common denominator agreeable to stakeholders of what would constitute a future-proof and robust governance framework.

For example, few model developers seem to fare well in effectively ensuring the safety of their models. The [2025 AI Safety Index](#) by the Future of Life Institute, which grades the biggest players' safety and security measures, finds that "the industry is fundamentally unprepared for its own stated goals." Further, "only 3 of 7 firms report substantive testing for dangerous capabilities linked to large-scale risks such as bio- or cyber-terrorism." Project Midas's [Seoul Commitment Tracker](#) similarly grades companies on their respect for the Seoul Frontier AI Safety Commitments, finding only one company (Anthropic) faring relatively well (B-), while other leading frontier AI labs were graded C or worse. [SaferAI's rating](#) of risk management practices among companies that published their safety frameworks concludes that "all companies currently have weak to very weak risk management practices," with the best in class scoring 34 percent in risk maturity.

This bolsters arguments that lawmakers should move from voluntary frameworks to legally sanctioned regimes, especially if the goal is to ensure models are secure enough to prevent CBRN attacks, cyberattacks, and loss of control risks. Even still, moving to a legally sanctioned regime does not seem good enough as a long-term solution in ensuring models are secure against catastrophic risks. The [2026 International AI Safety Report](#) reinforces this concern: Even as risk management practices become more structured, the report warns that sophisticated attackers can often bypass current safeguards and that the real-world effectiveness of many protections remains uncertain—suggesting that legal mandates alone, however well designed, cannot bridge the gap between regulatory intent and actual protection against catastrophic harm.

The fact that most of the frameworks analyzed in this article mandate regular review of safety frameworks suggests that policymakers are attuned to the possibility of emergent risks arising from exploration of model capabilities, as well as a need to evolve governance approaches as the underlying technology develops, which is already acknowledged by both [researchers](#) and [industry](#). Furthermore, risk mitigation measures might need to adapt as the model and its deployment context evolve. These provisions incorporate some element of iteration into responsible frontier AI governance, though they are not a substitute for directly updating the underlying governance frameworks to account for significant technological developments, such as rapid advancements in autonomous agentic AI. Similarly, many frameworks include the need for developers to define risk thresholds. Defining these thresholds may pose a net benefit to developers by providing them with a clear picture of when the model performs in an unacceptable way and needs action.

At the same time, the analysis shows that some measures still have not found consensus across the states, such as a requirement that developers refrain from deploying a model if the risk it would introduce is too high. These types of provisions have likely been discarded due to opposition from industry and a concern that they are overly prescriptive and burdensome. The case of Anthropic’s Mythos—a model the company chose not to release publicly, instead making it available to a limited set of vetted organizations, over concerns about its security vulnerability identification capabilities—illustrates both the promise and the limits of voluntary self-governance. In declining to deploy Mythos, Anthropic acted consistently with the Seoul Frontier AI Safety Commitments, which include a pledge to withhold models deemed to pose unacceptable risks. Yet the adequacy of this outcome rests entirely on the disposition of the developer. No authority could have compelled that decision; a different company, whether a signatory to those commitments or not, might well have reached a different conclusion for commercial or strategic reasons. Only a legally binding regime can guarantee such restraint as a matter of right rather than goodwill. Notably, an earlier version of New York’s RAISE Act would have obligated developers to refrain from deploying high-risk models marketed in the state, but the provision was removed from the final text.

As cases like Mythos’s become less exceptional, and as the potential for catastrophic harm from frontier models grows harder to dismiss, the case for legally mandating deployment restraint is likely to become stronger. The most draconian measure, the possibility of enacting a full shutdown, has not made it into law, reflecting the fact that states are open to policy experimentation but are also responsive to demand signals from industry and other stakeholders.

Reporting of incidents has strong support across many policy initiatives as a valuable measure to ensure rapid response and enable comprehensive information sharing. There is, however, significant disagreement on the details of implementation. To maximize the effectiveness of such provisions, policymakers should define a sensible range of time frames to report incidents, possibly with a graduated approach according to the gravity of the harm.

Many of the frameworks do not include the appointment of senior personnel to oversee the safety framework, even though several U.S. state bills [make](#) a more general reference to “instituting internal governance practices to ensure implementation.” This may demonstrate a reluctance by policymakers to impose personal accountability by law, as well as pushback from an industry comfortable

with self-implementation of safety requirements but opposed to highly prescriptive governance requirements that delineate specific top-level managers responsible for the safety of models.

There is also substantial variability in the actors covered across the different frameworks. The approach in RAISE and S.B. 53, which both impose obligations only on model developers, is in line with the EU AI Act and likely more acceptable to industry. In contrast, Illinois S.B. 315 and Michigan H.B. 4668 also impose obligations on third-party evaluators, who often are independent researchers. These evaluators would be obliged to produce a report and identify issues of noncompliance by model developers. Evaluators also could face penalties if they do not comply with these obligations. Imposing obligations on other actors in the frontier AI value chain could disproportionately affect small businesses integrating AI models into their applications. Moreover, it may have a chilling impact on the nascent ecosystem of independent researchers performing external evaluations, stifling the formation of a competitive marketplace and diminishing the overall benefit of independent evaluations.

The United States and European Union also remain fundamentally distant in defining the scope of what frontier AI is and what specific risks it can introduce. In U.S. state bills, the scoping of covered models largely follows the compute threshold introduced by Executive Order 14110, with some attempts to cover distilled models such as DeepSeek-R1. To the compute threshold, California added a financial criterion related to the entity, rather than the model, ensuring that only the largest AI providers would be covered. By contrast, the European Union covers a wider range of models and foresees the possibility of directly designating smaller models as posing “systemic risk,” which could also apply to DeepSeek. A previous version of California S.B. 53 offered a similarly flexible possibility of direct designation that other U.S. lawmakers should consider integrating; this would help prevent covering models they did not intend to cover and make the legislation more future proof.

Lastly, as U.S. state bills evolve during the legislative process, they tend to shed their most controversial provisions while learning from one another, growing increasingly similar over time. Looking at the overall trajectory, the active bills analyzed in this article have moved away from overly prescriptive approaches—such as the one proposed by the now-defunct California S.B. 1047—through an intermediate stage still displaying a range of obligations for large frontier model developers, comparable to several international frameworks, before converging toward narrower disclosure requirements rather than substantive obligations. California S.B. 53 pioneered this approach, followed by Massachusetts S. 2630 (which has not yet been approved) and RAISE. This convergence may reflect a “lock-in effect” generated by the final, less controversial version of S.B. 53. Once that benchmark was established, the other two bills shifted rapidly in the same direction. The remaining Michigan bill may ultimately follow the same path as it advances through its legislative session.

This emerging template is also starting to influence the debate at the federal level. The bipartisan discussion draft of the [Great American AI Act](#), released by Representatives Jay Obernolte and Lori Trahan on June 4, 2026, structures its frontier AI governance title around the same core pillars that state bills had already established, while also incorporating the independent verification organization model [recently enacted](#) in Virginia and Connecticut. A federal legislative approach consistent with state-level approaches is perhaps the clearest sign yet of convergence in elements of frontier AI governance in the United States, strong enough to pair with a proposed three-year preemption of state AI development laws.

Conclusion

State initiatives to regulate frontier AI reflect the role that these jurisdictions have long played as experimental testing and vetting grounds for new policy approaches. Despite fears from the Trump administration and its congressional allies that state actions will significantly hinder AI development, the laws and bills on frontier AI safety that have passed public and stakeholder scrutiny are largely consistent with the administration's stated aims in the AI Action Plan to increase security against malicious actors, minimize cyberattacks, and stop the creation of unauthorized CBRN weapons and explosives. In that sense, these state laws are a strategic asset that the federal government can leverage in its attempts to develop a nationwide framework that can address public mistrust of AI and prevent state fragmentation.

The Trump administration's response to Mythos might also signal that some degree of structured government oversight of frontier models may now be considered increasingly necessary across the political spectrum. The convergence of federal and state instincts around safety, however tentative at the federal level, makes the case for state disclosure frameworks all the more defensible: not as obstacles to innovation but as a durable, democratically grounded expression of a shared underlying concern—one that would complement federal oversight efforts and incentivize virtuous behavior within industry. Yet predeployment approval, however valuable, addresses only one moment in a model's existence. As Aalok Mehta notes in a [2026 CSIS analysis](#), a licensing or pretesting regime “can only be part of the story”; effective AI governance also requires continuous postmarket monitoring, incident response, and enforcement mechanisms that adapt as model capabilities evolve. State disclosure frameworks built around ongoing transparency obligations rather than one-time certification are structurally well suited to contribute to a longer arc of accountability.

Ignoring the work of state legislators would indeed mean forgoing a valuable opportunity to build on what states have already accomplished. States are converging on a series of practices that could anchor U.S. AI security governance: requiring an SSP, setting clear transparency rules, assessing critical risks, and mitigating potential catastrophes. These measures already reflect best practices from industry leaders and foreign governments. Even some [previous critics](#) of U.S. state bills seem to view a framework such as California S.B. 53 as a sensible and rather light-touch intervention to ensure the safety of the most powerful models. The Great American AI Act discussion draft, which takes a frontier AI governance approach consistent with the work of state legislators, is further indication that state-level work is and should be informing a possible federal framework.

The U.S. government should not interpret these developments as the result of importing EU-style rules; rather, they are the outcome of a parallel effort happening in multiple jurisdictions at the same time and strong industry involvement. Drawing upon these common frameworks would spare the federal government from creating novel approaches from scratch, strengthen U.S. security against dangerous AI and adversarial attacks, and enhance the trustworthiness of leading models (and the associated tech stack) abroad. ■

Laura Caroli is a political adviser on the Internal Market and Consumer Protection Committee for the Socialists and Democrats Group in the European Parliament. She was previously a senior fellow with the Wadhwani AI Center at the Center for Strategic and International Studies (CSIS) in Washington, D.C. **Aalok Mehta** is the director of the Wadhwani AI Center at CSIS.

The authors wish to thank Elena Gurevich, David Evan Harris, Cobun Zweifel-Keegan, Alexandra Tsalidis, and Matt Mande for their support in reviewing this article.

The views expressed herein are solely those of the authors and do not represent the positions of the European Parliament or the S&D Group.

This report is made possible by general support to CSIS. No direct sponsorship contributed to this report.

This report is produced by the Center for Strategic and International Studies (CSIS), a private, tax-exempt institution focusing on international public policy issues. Its research is nonpartisan and nonproprietary. CSIS does not take specific policy positions. Accordingly, all views, positions, and conclusions expressed in this publication should be understood to be solely those of the author(s).

© 2026 by the Center for Strategic and International Studies. All rights reserved.