

THE STATE OF STATE AI

Legislative Approaches to AI in 2025

| OCTOBER 2025



AUTHORS

Justine Gluck

Policy Analyst, AI Policy and Legislation, Future of Privacy Forum

Beth Do

Policy Fellow, Future of Privacy Forum

Tatiana Rice

Senior Director of U.S. Legislation, Future of Privacy Forum

ACKNOWLEDGEMENTS

Thanks to Bridget Egan for her research contributions.



The Future of Privacy Forum (FPF) is a non-profit organization that serves as a catalyst for privacy leadership and scholarship, advancing principled data practices in support of emerging technologies. Learn more about FPF by visiting fpf.org.



All FPF materials that are released publicly are free to share and adapt with appropriate attribution. [Learn more.](https://fpf.org)

EXECUTIVE SUMMARY

The Future of Privacy Forum's (FPF) report, *The State of State AI: Legislative Approaches to AI in 2025*, identifies the key trends in private-sector artificial intelligence (AI) policymaking reflected in major state bills enrolled, enacted, or advanced in 2025. The report identifies five key takeaways—

1. **State lawmakers moved away from sweeping frameworks regulating AI**, towards narrower, transparency-driven approaches.
2. **Three key approaches to private sector AI regulation emerged**: use and context-specific regulations targeting sensitive applications, technology-specific regulations, and a liability and accountability approach that utilizes, clarifies, or modifies existing liability regimes' application to AI.
3. **The most commonly enrolled or enacted frameworks** include AI's application in healthcare, chatbots, and innovation safeguards.
4. **Legislatures signaled an interest in balancing consumer protection with support for AI growth**, including testing novel innovation-forward mechanisms, such as sandboxes and liability defenses.
5. **Looking ahead to 2026**, issues like definitional uncertainty remain persistent while newer trends around topics like agentic AI and algorithmic pricing are starting to emerge.

TABLE OF CONTENTS

Executive Summary	1
I. Introduction	3
II. Four Thematic Approaches Characterize the 2025 State AI Legislative Landscape	4
III. Overview of Approaches to Private-Sector AI Regulations	6
A. Use or Context-Specific Approaches to AI Regulation	6
B. Technology-Specific Approaches to AI Regulation	9
C. Liability and Accountability Approaches to AI Regulation	12
IV. Defining the Next Wave: State AI Trends Heading Into 2026	15
A. Definitional Uncertainty Remains a Core Challenge	15
B. The Emerging Governance Question Around Agentic AI	16
C. Algorithmic Pricing as a New Focus for Consumer Protection	16
V. Conclusion	17

INTRODUCTION

In 2025, state legislatures across the United States accelerated their focus on artificial intelligence (AI), proposing a wide range of regulatory frameworks across technologies and sectors. FPF tracked **210 bills introduced in 42 states** that could directly or indirectly affect private-sector AI development and deployment. By concentrating on industry-facing legislation (rather than hundreds of additional bills that only reference AI in passing, update criminal codes, support workforce training or establish task forces), this report highlights the measures most likely to create compliance implications for companies developing or deploying AI systems. Only eight states did not introduce any bills meeting this threshold, underscoring the nationwide interest in AI.

With federal AI legislation yet to advance in Congress, and debates over an [AI state moratorium](#) marking tensions between federal and state roles in AI regulation, state legislators moved quickly to advance state AI bills. Despite this legislative activity, few proposals ultimately became law. Around **9%** of the bills tracked by FPF were enrolled or enacted, with most of those bills focused on government use of AI or state investment strategies, rather than imposing direct obligations on industry.¹ Substantive compliance frameworks—such as those targeting high-risk automated decision-making technologies (ADMTs)—faced hurdles in 2025, as lawmakers shifted away from broad, framework-style laws like [Colorado's 2024 AI Act](#) toward narrower, disclosure-based measures tailored to specific use cases and technologies.

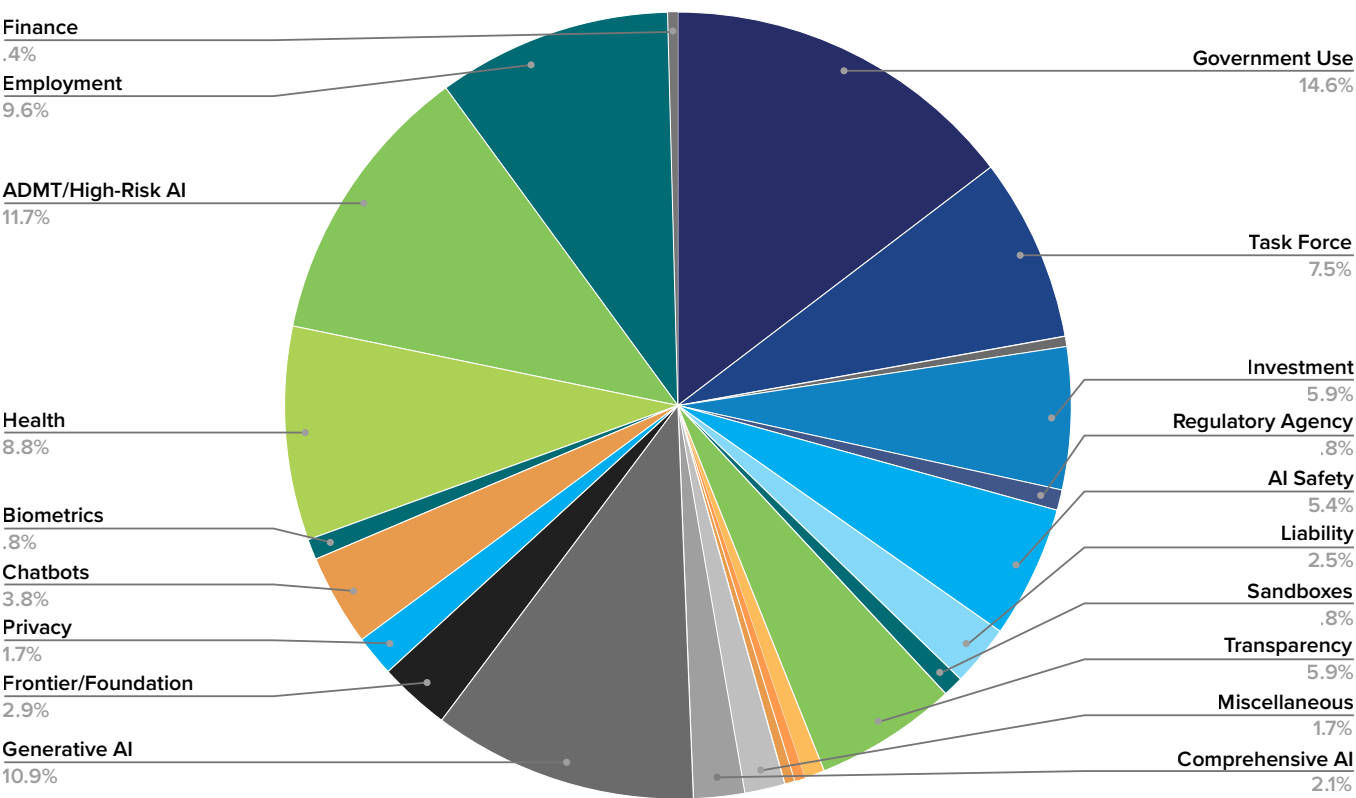
This report examines these trends by grouping state legislation into three primary approaches to private-sector AI regulation: use- and context-specific measures, technology-specific measures, and liability and accountability frameworks. By examining these trends, this report offers insight into how state lawmakers thought about AI in 2025 and the future of legislative trends in 2026. Together, these sections provide an overview of the major trends; insights into specific regulatory approaches; and a forward look at emerging themes that may shape legislation in 2026.

II. Four Thematic Approaches Characterize the 2025 State AI Legislative Landscape

As of the date of this report, **210 AI-related bills** were introduced in U.S. state legislatures that could directly or indirectly affect private-sector AI development and use. Of the bills directly affecting the private-sector, **eleven have been signed into law and nine are awaiting Executive action.**²

While other legislative trackers estimate that over 1,000 AI-related bills were introduced this year,³ this report applies a more targeted methodology, focusing attention on measures likely to meaningfully influence industry practices and the compliance landscape. This report excludes: bills and resolutions that merely reference AI in passing; updates to criminal statutes; and legislation focused on areas like elections, housing, agriculture, state investments in workforce development, and public education (which are less likely to involve direct obligations for companies developing or deploying AI technologies). This report does include bills regulating government agencies that are likely to impose requirements for private sector vendors.

Chart 1: State AI Bills by Category

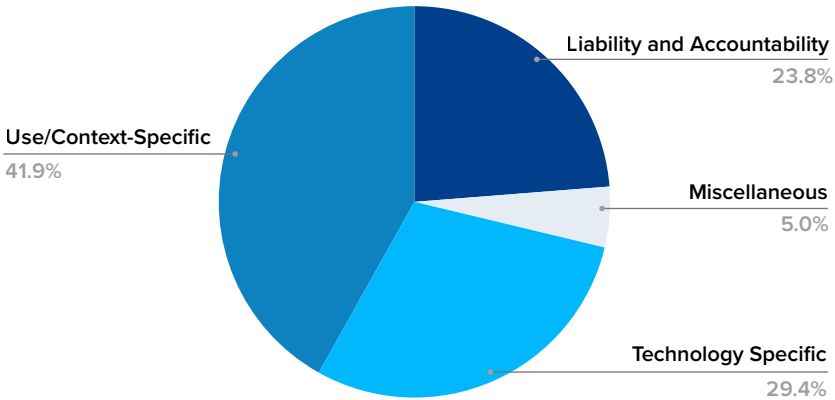


In 2025, AI-related legislation expanded not only in volume but also in diversity of approaches, with at least **eighteen** distinct types of AI-related bills introduced (see *Chart 1 above*). To bring greater clarity to this rapidly-evolving policy landscape, FPF developed a classification framework comprising four overarching thematic approaches—

Use / Context-Specific Bills	Focuses on certain uses of AI in high-risk decisionmaking or contexts—such as healthcare, employment, and finance—as well as broader proposals that address AI systems used in a variety of consequential decisionmaking contexts. These bills typically focus on applications where AI may significantly impact individuals’ rights, access to services, or economic opportunities.
Technology-Specific Bills	Focuses on specific types of AI technologies , such as generative AI, frontier/foundation models, and chatbots. These bills often tailor requirements to the functionality, capabilities, or use patterns of each system type.
Bills Focused on Liability and Accountability	Focuses on defining, clarifying, or qualifying legal responsibility for use and development of AI systems, such as establishing liability standards, creating affirmative defenses, or authorizing regulatory sandboxes. These aim to support accountability, responsible innovation, and greater legal clarity.
Government Use and Strategy Bills	Focuses on requirements for government agencies’ use of AI that have downstream or indirect effects on the private sector, such as creating standards and requirements for agencies procuring AI systems from private sector vendors.

A more detailed analysis of the bills and their categorization can be found in **Supplement Table 2**, and each enrolled or enacted laws’ categorization is in **Supplement Table 3**. While many bills touch on multiple themes, this framework is designed to capture each bill’s primary focus and enable consistent comparisons across jurisdictions.

Chart 2: State AI Bills with Industry Obligations by Broad Category



Of the **210 AI-related bills tracked by FPF in 2025, approximately 30 percent of those introduced focused on government use of AI**. Among the remainder of bills, which directly regulated industry stakeholders, efforts were broadly distributed across issue areas. Over **40 percent** of these remaining bills fell within the “use or context-specific” category—targeting sectors like health—with 15 percent of those bills enrolled or enacted. The next most common category was “technology-specific” bills, addressing particular systems such as chatbots, generative AI, and foundation models. Approximately

29 percent of the bills tracked by FPF with direct industry obligations fell into this category, with 24 percent of those bills enrolled or enacted. Finally, **24 percent** of bills focused on “liability and accountability” issues, addressing the legal responsibility for AI systems while simultaneously considering tools for responsible innovation; 26 percent of these liability-focused bills were enrolled or enacted. Chart 2 highlights this distribution, showing the share of bills in each category while excluding government use and strategy bills. Bills in the “miscellaneous” category are primarily comprehensive AI legislation.

Partisan Representation in State AI Policy

AI remains an area of interest across both major political parties. **Tables 9 and 10** in the supplement show that while Democratic legislators **introduced more than 75%** of all AI-related bills, nearly **41% of the AI bills signed into law** were introduced by Republicans. This share is lower than earlier in the session, when, prior to California's large wave of enrolled or enacted bills, closer to 75% of enacted measures had been Republican-led. In particular, Republican-led bills often focus on liability protections and government use, while Democrat-led bills have tended to prioritize transparency and consumer protections. This split suggests that while interest in AI is bipartisan, the types of measures introduced and enrolled or enacted tend to differ by political party.

III. Overview of Approaches to Private-Sector AI Regulations

The following sections examine the three primary approaches to regulating private-sector AI development and use in 2025: **Use or Context-Specific Regulation**, **Technology-Specific Regulation**, and **Liability & Accountability**.⁴ While the broader **Overview** highlights general legislative trends, this section offers a closer look at how state lawmakers are translating those trends into concrete policy. Each approach is illustrated through analysis of both enrolled or enacted laws and key bills that advanced through at least one legislative chamber, highlighting key trends, legislative language, and recurring policy approaches. These observations shed light on how states were shaping obligations for private-sector AI actors in 2025. The subsequent sections analyze each approach in detail.

A. Use or Context-Specific Approaches to AI Regulation

In 2025, over **45%** of the **twenty enrolled or enacted** private-sector AI laws and key bills directly impacting industry focused on particular uses or contexts of AI deployment, while such use-based approaches also made up over **40%** of all AI-related bills introduced.⁵ These approaches targeted applications where AI systems make or meaningfully influence decisions that have legal, financial, health-related, or other substantial effects on individuals' lives and opportunities.

This approach reflects the view that AI risks stem more from how systems are used (especially in high-impact contexts) than from their underlying technology. As a result, this sort of state legislation often targets the responsibilities of deployers, who are the entities implementing the AI system directly with individuals or consumers.

Based on the **nine laws enrolled or enacted** and the **six additional bills**⁶ that passed at least one chamber in 2025 seeking to regulate AI based on its use or context, common themes across these measures include—

- › **Focus on health-related AI applications:** Legislatures concentrated on AI in sensitive health contexts, especially mental health and companion chatbots, often requiring disclosure obligations.
- › **High-risk frameworks arose only through amendments to existing law:** No new standalone “high-risk” or ADMT frameworks were enacted in 2025, as substantive obligations emerged only through amendments to existing laws.
- › **Growing emphasis on disclosures:** User-facing disclosures became the most common safeguard.
- › **Shift toward fewer governance requirements:** Compared to 2024 proposals, 2025 legislation shifted away from compliance mandates, like impact assessments, in favor of transparency measures.

Targeted Focus on Health Applications of AI:

Four of the enrolled or enacted laws in 2025 directly focus on specific aspects of healthcare-related AI, while two others—Connecticut’s [SB 1295](#) (enacted) and Utah’s [SB 149](#) (enacted)—adopt broader frameworks that also apply to AI systems used in healthcare settings. These health-specific laws primarily focus on limiting or guiding AI use by licensed professionals, particularly in mental health contexts. Looking beyond enrolled or enacted measures, nearly **9%** of all introduced AI-related bills tracked by FPF in 2025 focused specifically on healthcare.⁷ From a compliance perspective, most prohibit AI from independently diagnosing patients, making treatment decisions, or replacing human providers, and many impose disclosure obligations when AI is used in patient communications.⁸

For example, Illinois’ [HB 1806](#) (enacted) offers one of the more detailed compliance frameworks in 2025. The law bars licensed therapy professionals from using AI for anything beyond defined “supplementary support” functions. This requirement creates new compliance obligations for providers to carefully track how AI is deployed in clinical settings and ensure consent is obtained whenever required. While the statute does not mandate formal documentation practices, providers will likely need to maintain records of consent and AI use as a practical safeguard against enforcement risk. Additionally, the law’s general prohibition against offering AI therapy services without a licensed professional may require developers and providers to implement guardrails to prevent general-purpose models for such purposes.⁹

Enrolled or Enacted “High-Risk” AI Frameworks Solely Arose from Amendments to Existing Law:

In contrast to 2024, when [Colorado enacted](#) the Colorado Artificial Intelligence Act ([CAIA](#)),¹⁰ an AI law regulating different forms of “high-risk” systems used in consequential decisionmaking, no similarly broad legislation was passed in 2025. High-risk AI and ADMT-focused bills made up 11 percent of the 210 AI-related bills introduced in 2025, with only 5 advancing beyond one chamber and 19 failing to move past introduction.

Several jurisdictions advanced “high risk” approaches through amendments to existing laws or rulemaking efforts—many of which predate Colorado’s AI law but reflect a similar focus on automated decision-making systems across consequential decisionmaking contexts. These include the California Privacy Protection Agency’s (CPPA) [regulations](#) on ADMT, Connecticut’s [SB 1295](#) (enacted), and New Jersey’s [ongoing rulemaking](#)—all of which address automated systems that produce legal or similarly significant effects on consumers. Additionally, Utah’s [SB 226](#) (enacted) amended its existing generative AI law to regulate only “high-risk” consumer-facing interactions—specifically, those that could reasonably be relied upon for significant decisions related to financial, legal, or medical services.

A New Notice and Choice Regime?

Notice and choice has long been a cornerstone of data privacy, dating back to the 1973 Fair Information Practice Principles ([FIPPs](#)). While foundational, this approach has faced [criticism from advocates](#) for placing burdens on consumers and from industry for hindering benign or beneficial data uses absent consent.

A similar model is now re-emerging in AI legislation. Though not new—early versions appeared in the Fair Credit Reporting Act—recent state AI laws increasingly require disclosures when individuals interact with AI systems, aiming to inform users and, in some cases, support enforcement of existing rights under other laws. Consumer and labor [advocates](#) argue that making the use of AI visible better equips consumers and workers to utilize existing rights under civil rights, consumer protection, and privacy laws.

Whether this trend will enhance transparency or replicate the shortcomings of privacy notice regimes remains to be seen, but it reflects a renewed focus on user agency in the age of AI.

Greater Focus on Disclosures to Individuals:

Eight of the enrolled or enacted laws and regulations require that individuals be informed when they are interacting with, or subject to decisions made by, an AI system.¹¹ The nature of the required disclosures varies: more targeted laws tend to mandate straightforward notices, while broader frameworks—such as those integrated into state consumer data privacy laws—often require more detailed and substantive disclosures about the use, purpose, and data use of the AI system. For instance, the chart below highlights how Utah and New Jersey have taken contrasting approaches to regulating high-risk AI (one more targeted, the other more comprehensive), which is similarly reflected in the differences in their disclosure requirements.

Utah’s SB 452	New Jersey’s Data Privacy Act Draft Regulations
<i>A targeted law concerning mental health chatbots</i>	<i>A data privacy law with provisions broadly regulating ADMT</i>
Requires suppliers to “clearly and conspicuously” notify users before an interaction that they are interacting with an AI technology.	Requires controllers to disclose to consumers the categories of data used for the automated profiling, types of decisions made, any evaluations for accuracy, fairness, or bias, and opt-out procedures.

Trend Toward Fewer Substantive Governance Requirements:

Compared to prior years, fewer of the enrolled or enacted laws utilizing this approach imposed substantive obligations on businesses, such as requirements to conduct impact assessments or maintain risk management policies. While some broader frameworks, particularly those arising as amendments to state data privacy laws, retained these obligations, most enrolled or enacted AI-specific bills did not. For the few laws that did include governance-related processes, the obligations were generally “softer,” such as being tied to an affirmative defense (e.g., Utah’s [HB 452](#) (enacted), which offers an affirmative defense to chatbot suppliers who maintain a governance policy) or satisfied through adherence to federal requirements (e.g., Montana’s [SB 212](#) (enacted), which requires a risk management policy for critical infrastructure facilities but allows a plan prepared under federal requirements to constitute compliance). Many other key bills that initially included governance requirements later removed or narrowed them to secure passage, such as Connecticut’s [SB 2](#) (proposed), which began as a broad “high-risk” framework but was ultimately passed by the Senate in a substantially pared-back version focused primarily on transparency (the bill was not ultimately enacted).

A Shifting Landscape in 2025

The political environment surrounding AI policy in the United States shifted notably in 2025, shaping the types of state legislation that gained momentum. Broad, comprehensive AI frameworks, such as the [Colorado AI Act](#), faced opposition, particularly from industry and free market advocates who warned such laws could [stifle innovation](#) and harm U.S. competitiveness. Narrow sector or use-based regulations gained traction. These targeted approaches drew bipartisan support across red and blue states.

- › **Rejection of the EU AI Act:** Parallel debates also surfaced around rejecting an “EU-style” approach, with critics expressing [strong concern](#) that replicating the EU AI Act could create regulatory overreach. However, civil society groups [argue](#) that even comprehensive frameworks like the Colorado AI Act or Virginia’s [HB 2094](#) (vetoed) are significantly narrower in scope and that highlighting that superficial similarities with the EU AI Act, such as shared definitions or structural references, mask substantial differences in substance. Notably, while broad “high-risk” frameworks drew criticism for resembling EU-style regulation, several Republican-led states adopted elements akin to the EU AI Act’s “prohibited AI practices” within government-use laws like Texas’ [HB 149](#) (enacted) and Montana’s [HB 178](#) (enacted), suggesting a growing distinction between restrictions that may be more acceptable in the public sector than for private industry.
- › Colorado’s definitional framework remains quietly influential: While no states enacted legislation akin to the Colorado AI Act in 2025, its definitional framework, particularly the focus on “high-risk” systems and “consequential decisions,” continued to influence legislation across party lines, including Utah’s [SB 226](#) (enacted) and Montana’s [SB 212](#) (enacted). Additionally, some enacted bills regulating government use of AI followed this structure including Arkansas’ [HB 1958](#), Kentucky’s [SB 4](#), and Texas’ [SB 1964](#).

B. Technology-Specific Approaches to AI Regulation

In 2025, state lawmakers increasingly introduced bills targeting specific *types* of AI technologies, rather than just their use contexts. Although technology-specific bills made up about **19%** of AI legislation **introduced** in 2025, they accounted for a disproportionate share of enrolled or enacted laws. Of the **twenty** state AI laws enrolled or enacted in 2025 that set direct obligations for industry, **five** of those laws specially address chatbots, pointing to the topic’s growing legislative attention.¹²

This report highlights **fifteen** notable technology-specific measures (**ten** enrolled or enacted, **five** advancing at least one chamber).¹³ Across these bills, three themes stand out—

- › **Chatbots as a key legislative focus:** Several new laws focused on chatbots, particularly “companion” and mental health chatbots, introducing compliance requirements for user disclosure, safety protocols to address risks like suicide and self-harm, and restrictions on data use and advertising.
- › **Frontier/foundation models regulation reintroduced:** California and New York revived foundation model legislation, building on 2024’s California’s [SB 1047](#) but with narrower scope and streamlined requirements. Similar bills surfaced in [Rhode Island](#), [Michigan](#), and [Illinois](#), signaling sustained interest in overseeing the most powerful AI systems.

- **Generative AI proposals centered on labeling:** A majority of generative AI bills in 2025 focused on content labeling—either user-facing disclosures or technical provenance tools like watermarking—to address risks of deception and misinformation. These measures signal legislative interest in transparency and consumer protection.

Targeted Focus on Chatbots, Particularly for Mental Health and Companions:

In 2025, chatbots—particularly those marketed as “companion chatbots” or chatbots deployed in mental health contexts—drew heightened legislative attention as lawmakers responded to recent [court cases](#) and [high-profile incidents](#) involving chatbot interaction. **Seven** chatbot-specific bills advanced at least one chamber, with **five** enrolled or enacted, compared to none signed in 2024. Most measures fell into two overlapping themes: (1) user identity disclosure and notification; and (2) safety protocols for emotionally sensitive contexts—

- **Disclosure Requirements as the Central Theme of 2025 Chatbot Bills:** Six of the seven key chatbot bills include a requirement for chatbot operators or suppliers to notify users that the chatbot is *not* human. Each bill mandates that the notification be “clear and conspicuous,” though they differ in how prescriptive they are about the timing, format, and language of the disclosure.¹⁴ California’s [SB 243](#) (enrolled), for instance, would require reminders that the chatbot is artificially-generated every three hours of use for *minor* users. These measures reflect a trend toward ensuring user awareness and transparency, while avoiding other compliance obligations like audits or risk assessments.
- **Chatbot Safety Protocols Emphasizing Suicide Risk and Self-Harm:** Several bills also introduced safety-focused provisions, particularly around suicide risk and self-harm. New York’s [S-3008C](#) (enacted), for example, prohibits offering AI companions without protocols that takes “reasonable efforts” to detect suicidal ideation and direct users to crisis resources. These requirements signal growing legislative concern about manipulative chatbot design and follow high-profile incidents in which chatbots allegedly encouraged self-harm.
- **Other Chatbot Accountability Measures Beyond Disclosure and Safety:** Beyond disclosure and safety, some bills experimented with accountability measures tied to privacy and advertising. Utah’s [SB 452](#) (enacted), for example, prohibits mental health chatbots from promoting products during conversations unless clearly labeled as advertising. California’s [AB 1064](#) (enrolled) initially highlighted concerns about how personal data from chatbot interactions with youth may be collected or used, though this provision was ultimately amended out of the final bill.

The Legal Fallout of Companion Chatbots

In the past year, multiple lawsuits have emerged regarding minors using “companion chatbots,” chatbots designed to simulate empathetic conversations and adapt to users’ emotional needs.

In both [Utah](#) and [Florida](#), the state AGs filed a complaint against Snapchat and its MyAI chatbot. Both states are bringing claims of potentially deceptive and exploitative treatment of minors and insufficient data-collection notices. A parent in Florida [sued](#) Character.AI (C.AI), a companion AI chatbot service, for negligence and the wrongful death of her teenage son, who took his own life after the chatbot allegedly encouraged him to do so.

Similarly, two parents are [suing](#) C.AI in Texas, after the bot engaged in inappropriate and explicit conversations with their children. Both suits allege that C.AI lacked safety measures, advertised an unsafe product to children, and employed an addictive design model, unjustly enriching themselves with the data of minors.

Frontier/Foundation Model Legislation Reintroduced:

In 2025, lawmakers introduced frontier/foundation model legislation centered on preventing “catastrophic risks” from the most powerful AI systems, like large-scale security failures that could lead to human injury or harms to critical infrastructure. With consideration to these types of threats, policymakers set these bills apart from ADMT or high-risk AI legislation that targets discriminatory outcomes and individual harms in specific domains. **Two key bills** that are both awaiting Governor signature, the California Transparency in Frontier Artificial Intelligence Act ([TFAIA](#) or SB 53) and the New York Responsible AI Safety and Education Act ([RAISE Act](#)), encapsulate this approach. These laws would regulate “large developers” ([RAISE Act](#)) or “frontier developers” ([TFAIA](#)), which are defined as developers employing high computing power thresholds, more than 10^{26} integer operations, in addition to cost qualifiers. The laws are scoped to apply in use cases where deployment could result in “catastrophic risk” ([TFAIA](#)) or “critical harm” ([RAISE Act](#)).

Key themes from this approach include:

- **Focus on Developers and Substantive Governance:** Unlike most use or context-specific bills that emphasize disclosures by deployers, the foundation model bills focus primarily on creating substantive safety governance requirements for model developers.
- **Streamlined Safety Requirements:** The foundation model bills of 2025 focus on instituting risk-mitigation measures, such as requiring written safety and security protocols that include information such as internal controls and mitigation steps. California’s [TFAIA](#) goes further by requiring a public-facing transparency report outlining any internal or third-party risk assessments, as well as protections to employee whistleblowers. However, compared to 2024, the 2025 bills notably streamline compliance requirements, which may reflect a growing legislative preference to avoid overly complex technical mandates that might inhibit innovation. For example, both bills omit the earlier version’s requirement for third-party audits, as well as for “full model shutdown” capabilities, which many critiqued as being technically challenging and inhibiting open-source development.¹⁵

After the Veto: California’s Frontier AI Report

In 2024, California State Senator Wiener (D) introduced [SB 1047](#), which aimed to establish safety and oversight requirements for developers of frontier AI models. The bill passed both chambers of the legislature but was ultimately vetoed by Governor Newsom (D). Following the veto, Governor Newsom convened the [Joint California Policy Working Group on AI Frontier Models](#) and directed the working group to develop a comprehensive [report](#) on state-level governance.

Key takeaways of the report include: the importance of incorporating early design choices in building flexible and robust policy frameworks, aligning incentives to leading safety practices, implementing whistleblower protections and third-party evaluations to increase transparency, and adjusting policy intervention thresholds to AI governance goals.

The Majority of Generative AI Bills Focused on Content Labeling, Including User Disclosures and the Tagging of Provenance Data:

Amid growing concerns about misinformation and consumer deception, lawmakers introduced legislation in 2025 addressing labeling for generative content. The majority of 2025 generative AI bills focused on content labeling, user disclosures, and provenance data.¹⁶ These proposals generally targeted providers or operators of AI systems that generate text, images, or other media for public use, rather than establishing broad governance frameworks. Most bills addressed specific concerns such as real-time

warnings to users or traceability of outputs, with additional proposals exploring themes like training data transparency or content ownership.

Key themes of this approach include:

- › **Labeling Generative Content:** In order to ensure consumers can distinguish AI-generated from human-created content, state lawmakers focused generative AI bills on content labeling, either required disclosures visible to users at the time of interaction, or a more technical effort of tagging of provenance or training data to enhance content traceability. Consumer-facing disclosure requirements appeared in several 2025 bills,¹⁷ while others—such as California’s [AB 853](#) (enrolled) and New York’s [S 6954](#) (proposed)—pursued provenance data or watermarking to improve traceability. These two content labeling approaches reflect complementary policy aims: one emphasizing real-time user awareness and liability, the other longer-term detection and mitigation of synthetic content.
- › **Consumer Warnings:** Like many chatbot bills, several generative AI bills focused on real-time warnings to consumers about the potential misuse of generative systems. For example, California’s [SB 11](#) (enrolled) would require providers of AI systems designed to create digital replicas to provide a consumer warning that unlawful use of the AI system to “depict another person without prior consent” may result in civil/criminal liability for the user, while New York’s [S 934](#) (proposed) would similarly require providers to post “clear and conspicuous” notices warning users of a generative AI system’s outputs’ possible inaccuracy.

C. Liability and Accountability Approaches to AI Regulation

In 2025, a newer approach emerged that focused on defining, clarifying, or qualifying legal responsibility for use and deployment of AI systems. Though most experts agree that existing law applies to AI, there is a notable gap in how that looks in practice. Therefore, unlike the use- or technology-specific approaches that create new independent laws governing AI, this approach focuses on using or refining existing legal tools. While some liability approaches looked to create new liability mirroring tort regimes for when an AI system caused harm, others looked to limiting liability based on roles in the AI value chain or in order to foster greater innovation. This past year, **eight** laws were enrolled or enacted and **nine** notable bills advanced at least one chamber in this category.¹⁸

Across these measures, state legislatures tested different ways to balance liability, safety, and innovation. Common themes include—

- › **New and clarifying liability regimes for accountability:** States tested both new liability frameworks and made clarifications to existing law, employing provisions like affirmative defenses to incentivize responsible practices and updating privacy and tort statutes to address AI-specific risks.
- › **Prioritization of innovation-focused measures:** States experimented with regulatory sandboxes that allow controlled AI development and some legislation introduced “right to compute” provisions to protect AI development and deployment.
- › **Enforcement tools and defense strategies:** Legislatures expanded Attorney General investigative powers (such as civil investigative demands) and introduced a variety of defense mechanisms, including specific protections for whistleblowers.

New and Clarifying Liability Regimes for Accountability:

Frameworks looking to create new liability regimes or clarify existing ones, took varied approaches based on their goal. While some aimed to create new liability regimes to encourage responsible AI practices and accountability, others aimed to encourage AI development by creating legal protections and greater regulatory clarity. Among these liability bills, the three general trends of approaches include—

- **Affirmative Defenses and Rebuttable Presumptions to Incentivize Responsible Practices:** Across numerous states, lawmakers looked to affirmative defenses, or legal claims that allow defendants to dismiss lawsuits based on certain grounds,¹⁹ as a solution for incentivizing responsible AI practices while maintaining flexibility and reducing legal risks for businesses. While not entirely new,²⁰ these provisions were observed in higher volume and across a variety of frameworks in 2025: from Utah’s mental health chatbot law ([HB 452](#), enacted) allowing an affirmative defense if a provider maintained certain AI governance measures to Texas’ [TRAIGA’s](#) (enacted) affirmative defense for entities that cure a violation and otherwise complies with recognized AI risk management frameworks. Alternatively, California’s [SB 813](#) (proposed) took a novel approach by allowing AI developers to use certified third-party audits as an affirmative defense in civil lawsuits.
- **Clarifying Existing Laws and Liability:** Other bills aimed to clarify existing law or liability to create clearer rules of the road and how AI does or does not fit into legal claims or regimes. In some cases, lawmakers adopted various amendments to existing laws to account for new AI risks or developments, such as [TRAIGA’s](#) amendment of the Texas biometric privacy law to account for AI training, while Connecticut updated its comprehensive data privacy law to account for AI. Some proposals looked more broadly, such as California’s [AB 316](#) (enrolled) which clarified that in tort claims, developers do not have a legal defense against claims for AI “autonomously” causing the harm. Arkansas’s [HB 1876](#) (enacted) similarly clarified ownership by granting rights in generated content to the person who provides input or data to a generative AI tool.

Prioritization of Innovation-Focused Measures:

As policymakers across the country seek to foster AI innovation, several legislatures have begun experimenting with novel regulatory approaches. Two key approaches include—

- **Sandboxes:** This year saw the enactment of new regulatory sandboxes in [Texas](#) and [Delaware](#), along with the [first official sandbox agreement](#) under Utah’s 2024 AI Policy Act ([SB 149](#)). Sandboxes allow participants to test emerging technologies within a controlled environment subject to government oversight, offering lawmakers a way to balance consumer protections against helping smaller businesses navigate legal risk. Sandbox provisions were also utilized in various other bills like Connecticut’s [SB 2](#).
- **Right to Compute:** This year a few states looked towards a novel concept, coined as the “[right to compute](#)” or the right to acquire and use AI technologies without government restriction. In enacting a “right to compute” Montana’s [SB 212](#) would prohibit the state from generally restricting the use or development of AI without a “compelling government interest.”
- **New Legal Protections to Support AI Innovation:** A few states have promoted mechanisms for new protections to support AI innovation and mitigate liability under various state laws and regulations. Texas’ [TRAIGA](#) ([HB 149](#), enacted) allows developers and deployers to maintain a *rebuttable presumption* of using reasonable care if they are compliant with relevant bill provisions. In Utah, [HB 452](#) (enacted), which includes a disclosure requirement for AI chatbots, allows for businesses to take advantage of an *affirmative defense* if they maintain certain AI governance measures. Other jurisdictions, such as California and New York, favor third party audits as part of an affirmative defense. California’s [SB 813](#) would allow AI developers to use certified third-party audits as a defense in civil lawsuits.

States Take Message from the Federal Administration

Several state legislatures are drawing inspiration from federal priorities as they shape their own AI frameworks. Provisions supporting innovation and legal clarity have begun to echo aspects of the White House's [America's AI Action Plan](#) released in July 2025, which emphasizes fostering innovation while reducing regulatory barriers.

For example, the Plan's Pillar One, "Accelerate AI Innovation," explicitly recommends establishing regulatory sandboxes: a trend reflected in Utah, Texas, and Delaware, where lawmakers are testing how controlled experimentation can balance oversight with innovation. State efforts to codify a "right to compute," such as Montana's [SB 212](#) (enacted), similarly align with the federal plan's calls to streamline bureaucratic hurdles and ensure broad access to AI technologies. These parallels show how some states are reflecting federal priorities to embrace a shared vision for innovation-friendly governance.

Emergence of a Variety of Enforcement Mechanisms and Defense Strategies:

States also experimented with a variety of enforcement mechanisms, including providing their Attorney Generals (AGs) with broad authority to investigate and enforce AI-related claims. Key themes include—

- › **Broad Investigative Authority Provided to State Attorneys General:** Several bills, including [TRAIGA](#) (enacted) and [Virginia HB 2094](#) (vetoed), provided for broad investigative authority to their state AGs to investigate and enforce AI-related claims. As part of their authority, these bills allowed AGs to issue a "civil investigative demand" (CID), a [discovery tool](#) used by AGs or agencies to obtain information (typically issued before a formal complaint is filed). While not a new legal tool, the information an AG can obtain may serve as an incentive for businesses to conduct certain practices and retain certain documentation. For example, under [TRAIGA](#), the AG may demand a broad amount of information about organizations' AI systems, including *any relevant documentation* reasonably necessary to conduct the investigation, including things not required under the substantive provisions of the law, such as information on data sources, model development processes, or safeguards implemented to mitigate risks.
- › **Whistleblower Protections:** In addition, states have proposed a variety of safety mechanisms for AI development and deployment to protect consumers, including specific protections for whistleblowers. For example, California's [SB 53](#) (enrolled) would protect whistleblower employees at large AI frontier model labs who report safety "critical risks" to the Attorney General. New York's [S 1169](#) (proposed) would similarly prevent developers and deployers of high-risk AI systems from restricting employees who disclose violations of the Act (which would set requirements for high-risk AI testing and transparency) to the Attorney General.

IV. Defining the Next Wave: State AI Trends Heading Into 2026

As states prepare for the 2026 legislative cycle, many are expected [to revisit](#) unresolved questions and familiar topics from 2025, while also confronting new themes. Definitions will likely remain a central challenge: although states refined AI terminology this year, variations continue to shape which systems fall under compliance—an issue likely to persist in 2026. As bills increasingly target specific capabilities and model types (e.g., frontier/foundation models), consistent definitions will be essential to reduce a patchwork of legislation. At the same time, early proposals around new applications of AI, such as agentic AI systems and algorithmic pricing, signal a new wave of policy debates.

This section highlights key areas to watch in 2026—

- › **Definitional uncertainty:** States continue to vary in how they define “artificial intelligence” itself, as well as definitions specific to frontier/foundation models, generative AI, and chatbots: raising questions about scope and consistency across jurisdictions.
- › **Agentic AI:** While still nascent in state legislation, AI “agents” capable of autonomous action are beginning to draw legislators’ attention, including early governance experiments such as sandboxes.
- › **Algorithmic pricing:** States are testing bills to regulate algorithmic pricing practices, focusing on discrimination, transparency, and consumer protection, a trend signaling its expansion in 2026.

The sections below explore these trends in greater detail, beginning with how states approached AI definitions in 2025 and turning to emerging issue areas, such as agentic AI and data-driven pricing, that may define the next era of state AI policy.

A. Definitional Uncertainty Remains a Core Challenge

As AI Regulation Becomes More Targeted, How States Define Key Systems Will Determine the Scope and Impact of New Laws:

In 2025, legislatures largely refined definitions across four categories—AI, frontier/foundation models, generative AI, and chatbots—but variations remain significant, especially among technology-specific terms (e.g. chatbots). These definitional differences already affect which technologies are regulated and will become more consequential as states increasingly enact AI legislation.²¹

- › **Definitions of AI are Broadly Built on the OECD Baseline:** Most states continued to base AI definitions on the Organization for Economic Cooperation and Development’s (OECD) [widely-adopted language](#), emphasizing a system’s ability to generate outputs that influence physical or virtual environments.²² For instance, Texas’ [HB 149](#) (enacted) adopts this definition verbatim, while New York’s [S-3008C](#) (enacted) expands upon it by specifying that AI systems “abstract such perceptions into models through analysis in an automated manner,” and then use inference to formulate options for action. This pattern reflects a broader trend of slightly adapting the OECD baseline to state-specific or legislation-specific contexts.
- › **Thresholds for Frontier/Foundation Model Definitions:** California’s Transparency in Frontier Artificial Intelligence Act (TFAIA, [SB 53](#), enrolled) and New York’s RAISE Act ([S 6453](#), enrolled) both introduced detailed, though slightly different, definitions for advanced AI systems. Each relies on a compute threshold of more than 10^{26} integer operations, but differ in scope for a cost threshold: [RAISE](#) focuses on a \$100 million “compute cost,” while [TFAIA](#) ties its standard to annual gross revenues in excess of \$500 million. [TFAIA](#) also defines a “foundation model” as one trained on a broad dataset, designed for general-purpose outputs, and adaptable across a wide range of tasks. Under this framework, “frontier models” are treated as a subcategory of foundation models. The [RAISE Act](#), by contrast, defines only “frontier models” and does not provide a definition for foundation models.

- **Generative AI Definitions Converge with Similar Broad Language, but Language Varies in Qualifiers and Scope:** Most 2025 bills converge on definitions for generative AI that describe models that emulate the “structure and characteristics” of training data to generate synthetic content across multiple formats. New York’s [S 6954](#) (proposed) and [A 6578](#) (proposed) describe generative AI as “self-supervised” models that generate “derived synthetic content,” while California’s [AB 853](#) (enrolled) uses nearly identical language but drops the “self-supervised” qualifier, slightly broadening its applicability. Other bills, like Massachusetts’ [H 90](#) (proposed), include even more concise definitions, defining generative AI as any system that generates content “based on the patterns of structures from its training data,” without specifying content types or self-supervision. Other states avoided defining “generative AI” directly, instead regulating adjacent terms like “deepfake” or “synthetic content.”²³
- **Definitions for Chatbots Remain the Most Fragmented, Though Common Language is Emerging for Companion Chatbots:** States used a wide range of terms to define chatbots in 2025 that carry meaningful implications for which systems are subject to regulation.²⁴ Still, repeated use of descriptors, like simulate, sustain, and human-like, suggests an emerging definitional baseline. Definitions of companion chatbots often highlight three traits: simulating human interaction, sustaining conversations, and addressing user well-being or social needs. Bills like New York’s [S-3008C](#) (enacted) and California’s [SB 243](#) (enrolled) and [AB 1064](#) (enrolled) include these three qualifiers. However, New York’s [S-3008C](#) definition extends that definition further by also including functions like retaining prior user interactions and asking “unprompted or unsolicited emotion-based questions.” Other 2025 bills define chatbots more broadly, focusing on users’ perception of the chatbot’s *humanness*, such as Maine’s [LD 1727](#) (enacted) which defines an “artificial intelligence chatbot” as any software that “simulates human conversation and interaction.” These definitions reflect a trend toward defining chatbots based on their potential to deceive users.

B. The Emerging Governance Question Around Agentic AI

The definitional disparities of 2025 highlighted lawmakers’ struggles to define fast-evolving technologies. Looking ahead, the rise of agentic AI underscores how those challenges may intensify as states confront even more advanced and autonomous systems. [Agentic AI](#) (or AI agents) refers to a type of AI program that is capable of autonomously understanding, planning, and executing tasks, moving beyond generative AI’s content creation and towards more complex functions. Although AI agents have many beneficial uses, their application also raises many of the same data protection [questions](#) raised by LLMs, such as [challenges](#) related to collecting and processing personal data for model training, security vulnerabilities, and explainability. Other emerging questions—such as how personalization shapes complex AI systems and what that means for user privacy and safety—will further complicate how states approach responsible deployment of agentic AI.

States are only beginning to test agentic AI’s potential in government use and beyond. In July 2025, Governor Youngkin launched the nation’s first agentic AI-powered “[regulatory reduction pilot](#)” to reduce regulatory burdens on state agencies, while Delaware enacted the country’s first regulatory sandbox dedicated to agentic AI ([H.J.R. 7](#)). Still, few bills have addressed agentic AI directly, with legislative focus largely centering on popularly adopted technologies, like chatbots. Notably, risk assessment frameworks common in current state AI laws may prove ill-suited for agents (as AI agents operate through multiple decision-making nodes, making sources of harm more difficult to trace and regulate), suggesting that governance approaches will need to adapt as these technologies become increasingly adopted.

C. Algorithmic Pricing as a New Focus for Consumer Protection

Along with definitional debates and early exploration of agentic AI, [lawmakers in 2025](#) also turned to algorithmic pricing and other practices involving algorithms and personal data. These proposals targeted practices in which large datasets and algorithms determine and adjust the prices or products offered to consumers. While dynamic pricing is not new, AI amplifies its scale, speed, and precision, raising concerns of discriminatory outcomes, reduced transparency, personal data misuse, and anti-competition.

2025 saw legislative activity on the state level focused on algorithmic pricing, such as New York's [S 3008](#) (enacted) that requires disclosure when “personalized algorithmic pricing” is used. Legislation in California ([AB 446](#), proposed) suggested prohibiting pricing based on personal information obtained through “electronic surveillance technology,” while another California bill ([SB 384](#), proposed) suggested prohibiting “price-setting algorithms” used by competitors in the same market if the algorithm uses “nonpublic input data.” Other states, including Colorado and Minnesota, introduced similar measures, signaling growing momentum.

Like agentic AI, algorithmic pricing illustrates how legislators are targeting emerging, fast-evolving applications where existing law may not anticipate AI's capabilities. In 2026, legislative activity on algorithmic pricing may grow, with reintroduced bills apt to adopt more precise definitions and stronger disclosure mandates or prohibitions. Lawmakers may drive debates on whether current consumer protection and anti-discrimination frameworks are sufficient to address the risks of algorithmic pricing, or whether this practice is a distinct challenge requiring new regulatory frameworks.

V. Conclusion

The 2025 legislative cycle shows a state-level regulatory environment that is highly active but still experimenting with approaches to AI governance. Legislatures explored a wide range of approaches, from disclosure requirements to liability frameworks to technology-specific rules, but few gained the consensus needed to become law. Of the **210 bills** tracked by FPF, only **20** (about **9%**) were enrolled or enacted. Instead, most measures reflected fragmented efforts aimed at transparency and accountability, often focused on particular sectors or on testing the edges of liability.

Looking ahead, definitional clarity and scope will remain central challenges, especially as lawmakers turn to more advanced AI models like frontier systems and emerging agentic AI. Early debates around these applications suggest that the breadth of policymaking will only continue to expand. Simultaneously, lawmakers are expected to return to key themes from the 2025 legislative cycle, including AI's application in healthcare, companion chatbots, and frontier/foundation models. As the 2026 legislative sessions open, state legislators are likely to revisit these unresolved issues from 2025. They may also confront a new wave of technologies and governance questions, ensuring that state-level AI regulation remains a rapidly evolving frontier.

ENDNOTES

- 1 See Supplement Table 1.
- 2 Upon publication of this report, bills in California and New York are still awaiting gubernatorial action. This total is limited to bills with direct implications for industry and excludes measures focused solely on government use of AI or those that only extend the effective date of prior legislation. See Supplement Table 1 for a list of the 20 enrolled or enacted bills.
- 3 See Supplement Tables 3 and 4; NCSL, “Artificial Intelligence 2025 Legislation,” July 2025, <https://www.ncsl.org/technology-and-communication/artificial-intelligence-2025-legislation>; Multistate.AI “Artificial Intelligence (AI) Legislation,” September 2025, multistate.ai/artificial-intelligence-ai-legislation; and BCLP, “U.S. State by State AI Legislation Snapshot,” September 2025, bclplaw.com/en-US/events-insights-news/us-state-by-state-artificial-intelligence-legislation-snapshot.html.
- 4 *Note:* Legislation focused exclusively on government use of AI is excluded from this section but may be referenced where relevant to broader policy trends.
- 5 See Supplement Table 12.
- 6 *Enrolled or enacted:* California’s [SB 7](#), California’s [AB 489](#), Connecticut’s [SB 1295](#), Illinois’ [HB 1806](#), Montana’s [SB 212](#), Nevada’s [AB 406](#), Texas’ [SB 1188](#), Utah’s [SB 226](#), and Utah’s [SB 452](#). *Advanced At Least One Chamber:* California’s [SB 238](#), California’s [SB 420](#), California’s [AB 1018](#), Connecticut’s [SB 2](#), New York’s [S 1169](#), and Virginia’s [HB 2094](#).
- 7 See pie chart above; employment bills were the highest sector introduced, comprising 10% of all introduced AI bills.
- 8 Examples of these bills include: California’s [AB 489](#) (enforcing title protections for health professionals against those developing/deploying AI systems); Illinois’ [HB 1806](#) (barring licensed therapists from using AI without consent); Texas’ [SB 1188](#) (requiring health care practitioners to review info obtained through AI); and Nevada’s [AB 406](#) (prohibiting unlicensed AI systems from providing mental healthcare, and restricting licensed providers’ use of AI in clinical care).
- 9 [Illinois’ HB 1806](#) Sec. 20
- 10 *Note:* In 2025, Colorado held a special legislative session during which Colorado [SB 4](#) was enacted, delaying implementation of the Colorado AI Act by five months, shifting from an effective date of February to June 30, 2026.
- 11 New York’s [S-3008C](#), Maine’s [LD 1727](#), Utah’s [SB 452](#), Nevada’s [AB 406](#), California’s [SB 7](#), California’s [AB 489](#), California’s [CCPA’s ADMT](#) rules, and New Jersey’s [Data Privacy Act draft regulations](#).
- 12 See Supplement Table 13.
- 13 *Enacted:* Arkansas’ [HB 1876](#), California’s [SB 243](#), California’s [SB 53](#), California’s [SB 11](#), California’s [AB 1064](#), California’s [AB 853](#), Maine’s [LD 1727](#), New York’s [S 6453](#), New York’s [S-3008C](#), and Utah’s [SB 452](#). *Advanced At Least One Chamber:* California’s [AB 410](#), New York’s [S 6954](#), New York’s [S 5668](#), New York’s [S 934](#), New York’s [A 6578](#).
- 14 New York’s [S-3008C](#) and [S 5668](#) require operators to notify users at the beginning of any chatbot interaction and at least once every three hours during ongoing conversations. Maine’s [LD 1727](#) requires that users be notified “in a clear and conspicuous manner” but does not dictate when the disclosure must occur. Utah’s [SB 452](#) includes a more user-responsive approach: it mandates disclosures before users access chatbot features and whenever the user asks whether AI is being used.
- 15 Fei-Fei Li, “The ‘Godmother of AI’ Says California’s Well-Intended AI Bill Will Harm the US Ecosystem,” Fortune, <https://fortune.com/2024/08/06/godmother-of-ai-says-californias-ai-bill-will-harm-us-ecosystem-tech-politics/?abc123>.
- 16 *Note:* Content labeling requires AI-generated material to be clearly identified as AI-produced. Provenance refers to attaching metadata that records how, when, and with what tools content was created, to help verify authenticity.
- 17 Bills include: Pennsylvania’s [HB 95](#), Massachusetts’ [HD 1222](#), Georgia’s [HB 478](#), and Illinois’ [SB 1792](#).
- 18 *Enrolled or enacted:* California’s [SB 53](#), California’s [AB 853](#), California’s [AB 489](#), California’s [AB 316](#), Montana’s [SB 212](#), New York’s [S 6453](#), Texas’ [HB 149](#), and Utah’s [HB 452](#). *Advanced At Least One Chamber:* California’s [AB 1405](#), California’s [SB 813](#), Connecticut’s [SB 2](#), New York’s [S 5668](#), New York’s [A 6578](#), New York’s [S 6954](#), New York’s [S 1169](#), and Virginia’s [HB 2094](#).
- 19 In Utah, [HB 452](#), which includes a disclosure requirement for AI chatbots, allows for businesses to take advantage of an *affirmative defense* if they maintain certain AI governance measures. Other jurisdictions, such as California ([SB 813](#)) and New York ([S 6453](#)), favor third party audits as part of an affirmative defense.
- 20 Colorado’s AI Act offers an affirmative defense for adherence to federal standards, like NIST frameworks.
- 21 See Supplement Tables 15, 16, 17, and 18.
- 22 The OECD defines AI as a “machine-based system that, for explicit or implicit objectives, infers, from the input it receives, how to generate outputs such as predictions, content, recommendations, or decisions that can influence physical or virtual environments. Different AI systems vary in their levels of autonomy and adaptiveness after deployment.”
- 23 Some bills, such as California’s [SB 11](#) or Arkansas’ [HB 1876](#), rely on broader “artificial intelligence” definitions while scoping provisions to capture generative AI use cases. Others define adjacent terms like “deepfake” or “synthetic content” rather than “generative AI” itself, effectively regulating the technology through related concepts.
- 24 Across legislation introduced in California, New York, Utah, and Maine, for example, lawmakers defined a range of terms: “AI companion,” “companion chatbot,” “mental health chatbot,” “artificial intelligence chatbot,” and “bot.”

