



U.S. State AI Legislation

How U.S. State Policymakers Are Approaching Artificial Intelligence Regulation

FPF U.S. Legislation Report

September 2024

Authored By: **Tatiana Rice**, Deputy Director, U.S. Legislation

Jordan Francis, Policy Counsel, U.S. Legislation

Keir Lamont, Director, U.S. Legislation



**FUTURE OF
PRIVACY
FORUM**

Executive Summary

As artificial intelligence (AI) becomes increasingly embedded in daily life, including critical sectors like healthcare and employment, state lawmakers have begun crafting regulatory strategies to address its heightened risks while recognizing its potential to unlock insights and enhance services. This report by the Future of Privacy Forum analyzes key trends and concepts from proposed and enacted U.S. state AI legislation. This Report highlights:

1. **State lawmakers are primarily focused on regulating AI used in consequential decisions** that significantly impact individuals' livelihood and life opportunities.
2. **A key goal for many lawmakers is to mitigate the risk of algorithmic discrimination**, either through prohibitions on AI systems with identified discriminatory risks or by establishing a duty of reasonable care to protect individuals from such discrimination.
3. **Most AI legislative frameworks create role-specific obligations**, including separate requirements for developers and deployers related to transparency, risk assessments, and AI governance programs.
4. **Common consumer rights around AI** include rights of notice and explanation, correction, and to appeal or opt-out of automated decisions.
5. **Alternatively, some lawmakers utilize a technology-specific approach** to address novel risks, such as those posed by generative AI or frontier or foundation models.

Table of Contents

Report

- I. Introduction**
- II. The ‘Governance of AI in Consequential Decisions’ Approach**
 - A. Scope: AI Systems Used in Consequential Decisions that Materially Impact Individuals
 - B. Provisions Regarding Algorithmic Discrimination
 - C. Common Obligations for Developers and Deployers
 - D. Common Consumer Rights
 - E. Investigation and Enforcement by the State Attorney General
- III. Alternative Technology-Specific Approaches**

Report Supplement

The Report is supported by a supplementary document containing materials that are referenced throughout the Report, including:

- **Table of Key Relevant Bills and Laws**
- **Table of Key Terms and Definitions**
- **Example Business Obligations Language and Sub-Requirements**
- **Example Individual Rights Language**

I. Introduction

Artificial Intelligence (AI), as both a technical field and a societal concept, encompasses a broad spectrum of technologies and systems, some of which have been in use for decades. AI holds immense potential to unlock critical insights, enhance efficiency, and boost economic competitiveness. It also creates significant and undeniable risks, including harms arising from inaccuracy and the potential for an exponential increase in societal discrimination if biases are embedded in the AI. Given the expanding role of AI across all sectors, including in critical sectors like healthcare, employment, and finance, policymakers around the world face the imminent and critical challenge of balancing policies that support the advancement of AI with the need to implement effective strategies that mitigate against potential risks.

It is within the context of these opportunities and challenges that a new class of U.S. state legislation has emerged, proposed by lawmakers who seek to establish responsibilities for the safe, fair, and transparent use of AI systems, particularly in significant decision-making processes.¹ In response to this increased legislative focus, the Future of Privacy Forum (FPF) has conducted a comprehensive review and analysis of key bills introduced in 2023 and 2024.² FPF also directly engaged with state policymakers on numerous proposals and helped convene a multistate AI policymaker working group. Through these interactions, we have gained additional insight into the current AI legislative process as well as the considerations that have thus far influenced state policymakers' decisions.

This report delves into the trends across these legislative efforts, examines core questions and issues, and offers key considerations for policymakers as they navigate the complexities of AI policy. In Section II, we examine in detail the most frequently introduced state legislative framework, **'Governance of AI in Consequential Decisions.'** The framework is drafted to apply to a broad range of entities and industries, offering one of the most comprehensive governance approaches currently under consideration for mitigating specific AI risks across various proposals and laws. Section III of this Report provides greater details on alternative approaches focused on particular technologies, such as generative artificial intelligence and frontier or foundation models.

¹ See, e.g., [California SB 1047](#) (proposed) (Aug. 19, 2024) (stating the purpose to “advance[] the development and deployment of artificial intelligence that is safe, equitable, and sustainable”); [California AB 2930](#) (proposed) (Aug. 15, 2024) (“The Unruh Civil Rights Act provides that all persons within the jurisdiction of this state are free and equal and . . . are entitled to the full and equal accommodations, advantages, facilities, privileges, or services in all business establishments of every kind whatsoever.”).

² See [Report Supplement](#), Table 1.

II. The ‘Governance of AI in Consequential Decisions’ Approach

In the context of regulating artificial intelligence, the most prevalent approach embraced by U.S. state legislators is the regulation of AI systems or tools used in consequential decision-making contexts that significantly impact individuals' livelihood and life opportunities. Often, these legislative proposals focus on the operative terms “**high-risk artificial intelligence system**” or “**automated decisionmaking tool**.”³ Typically, the goal of this framework is to create incentives for fairness, transparency, and oversight and accountability processes,⁴ mitigating algorithmic bias in areas typically covered by civil rights law, such as education, housing, financial services, healthcare, and employment. The most common provisions in this framework, include:

- A. Scope
- B. Provisions to Address Algorithmic Discrimination
- C. Developer and Deployer Obligations
- D. Consumer Rights; and
- E. Enforcement

A. Scope: AI Systems Used in Consequential Decisions that Materially Impact Individuals

When deciding on the definitional scope of either “high-risk artificial intelligence system” or “automated decisionmaking tool,” lawmakers often follow a common framework that can be broken down into five parts: (1) the definition of “artificial intelligence”; (2) the relevant context in which the law should apply; (3) the impact and role of AI system on the decision; (4) the individuals or entities subject to regulation; and (5) any necessary carve-outs or exclusions. Each is discussed below, with example definitions provided in the [Report Supplement](#), Table 2.

(1) Definition of Artificial Intelligence: When defining and determining what will constitute covered AI systems, there is consensus globally, nationally, and among U.S. state lawmakers as

³ “**High-risk artificial intelligence system**” was used in at least seven proposals in 2024, including the [Colorado SB 24-205](#) (enacted) (2024), [Connecticut SB 2](#) (proposed) (2024), and [Virginia HB 747](#) (proposed) (2024). “**Automated decision[making] tool**” or “technology” was used in at least ten proposals in 2024, including [California AB 2930](#) (proposed) (2024), [Washington HB 1951](#) (proposed) (2024), and [Oklahoma HB 3835](#) (proposed) (2024). Some federal proposals and regulations also utilize similar but varying terms. The [Office of Management and Budget](#) describes these systems as “**Rights-Impacting Artificial Intelligence**,” whereas the bipartisan “[Artificial Intelligence Research, Innovation, and Accountability Act of 2023](#)” (proposed) describes them as “**High-Impact AI Systems**.”

⁴ See, “[Artificial Intelligence Regulations April 18th Forum with State Legislators](#),” (featuring Alaska Senator Hughes, Colorado Senator Rodriguez, Connecticut Senator Maroney, Georgia Representative Jones, Texas Representative Capriglione, Virginia Delegate Maldonado) (discussing bipartisan state-level efforts and coordination to establish common sense guardrails for advancing trustworthy artificial intelligence).

reflected by varying laws, proposals, and frameworks to use the definition of “artificial intelligence” set forth by the [Organisation for Economic Cooperation and Development \(OECD\)](#):

“An AI system is 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.”

(2) Context: Most U.S. state laws and proposals focus on the use of AI systems or tools in key sectors protected by U.S. civil rights laws, such as employment, education, housing, and financial services. These areas are not only protected under existing U.S. civil rights law, but are also crucial to individuals' livelihoods, meaning that automated systems used in these contexts present a higher risk of harm. Under legislation that takes this approach, including the [Colorado AI Act](#) (enacted) and [California AB 2930](#) (proposed), decisions in these areas are considered “consequential decisions.”⁵ Amongst proposals that rely on this approach, the scope of “consequential decisions” almost always includes use of AI that would impact:

**Alternative Approaches:
Sector-Specific:**

Instead of broadly addressing “consequential decisions” across various areas, some legislative proposals and laws concentrate on specific sectors where automated systems are used. This approach is most commonly seen in regulations focused on AI in employment, such as [New York City Local Law 144](#) and [Illinois HB 3773](#) (2024), or AI in healthcare, as exemplified by [Georgia HB 887](#) (2024).

- Education enrollment or an education opportunity;
- Employment or an employment opportunity;
- Housing;
- A financial or lending service;
- An essential government service;
- Healthcare services;
- Insurance; and
- Legal services

Some proposals go further or contain more prescriptive lists. For instance, [California AB 2930](#) includes both a larger list of areas in scope and more prescriptive lists of what decisions in each area entails, such as essential utilities (including electricity, heat, water, internet, telecommunications access, and transportation), criminal justice (including risk assessments for pretrial hearings, sentencing, and parole), adoption services, reproductive services, and voting.

Once the scope of key areas is determined, lawmakers additionally consider the particular decision at play within the context. Almost all proposals use language inspired by the [EU’s General Data Protection Regulation](#), requiring that decisions produce “**legal or similarly**

⁵ Other proposals utilize alternative terms, such as “important life opportunities.” See, e.g., [District of Columbia Stop Discrimination by Algorithms Act](#), B. 25-0144 (proposed) (Feb. 2, 2023).

significant effects” in an individual’s life regarding a particular index of decisions.⁶ Some proposals have a more limited list of covered decisions like “provision or denial of” while others have a more exhaustive list, including “cost of” and “access to”. The language in the [Colorado AI Act](#) and [California AB 2930](#) provides a useful comparison:

Colorado AI Act	California AB 2930
“Consequential Decision” means “a decision that has a material, legal, or similarly significant effect on the provision or denial to any consumer of, or the cost or terms of... ”	“Consequential decision” means “a decision or judgment that has a legal, material, or similarly significant effect on an individual’s life relating to access to government benefits or services, assignments of penalties by government, or the impact of, or the cost, terms, or availability of, any of the following...”

(3) Impact and Role of the AI System: The most debated and difficult factor for U.S. state lawmakers to decide on has been the impact and role the AI system must play in the decision-making process in order for it to be in scope of regulation. Because AI systems vary widely in their functions, from simple algorithms to complex autonomous systems, not all AI technologies have the same impact on decision-making processes. This variability is compounded by the ambiguity surrounding what constitutes a "significant" impact on individuals' lives, often requiring a case-by-case analysis based on specific facts and circumstances.

Considerations on Advertising:

Advertising presents unique regulatory challenges, and there is currently no consensus on whether it should be included in AI regulations. Some lawmakers explicitly include advertising or indirectly cover it by targeting AI systems that affect an individual’s **access to** key areas and life opportunities, while others prefer to keep their proposals narrowly focused on regulating automated systems used in formal decision-making processes.

Lawmakers often focus on three key terms regarding AI's role in the decision-making process: "**facilitating decision making**" (lowest threshold), "**substantial factor**" (median threshold), and "**controlling factor**" (highest threshold).

- **Facilitating decisionmaking:** The [California Privacy Protection Agency](#) initially utilized “facilitate human decisionmaking” as the operative term to describe the role an automated system must play in decisionmaking process to bring it under the scope of the California Consumer Privacy Act (CCPA) opt-out rights.⁷

⁶ This language is also common in state comprehensive privacy laws in defining profiling decisions subject to opt-out rights. *E.g.*, [Virginia Consumer Data Protection Act](#), Code § 59.1-575; [Colorado Privacy Act](#), Rev. Stat. § 6-1-1303(10); [Connecticut Consumer Data Privacy and Online Monitoring Act](#), Gen. Stat. § 42-515(12).

⁷ The [latest version](#) of these draft regulations, this has been revised to a “**substantially facilitate**” standard.

- **Substantial factor:** This year, Connecticut Senate Bill 2 debuted “substantial factor.”⁹ The term was later carried into the Colorado AI Act, which ultimately defined “substantial factor” to mean when content generated by an AI system assists in making a consequential decision and is capable of altering the outcome.
- **Controlling factor:** The California legislature first utilized “controlling factor” in California AB 331 (2023) but left the term undefined.

(4) Regulated Entities: The ‘Governance of AI in Consequential Decisions’ approach typically accounts for **role-specific responsibilities and capabilities** in the AI system lifecycle, including the distinct (but occasionally overlapping) roles played by developers and deployers of AI systems. Developers, who build AI systems, and deployers, who use such systems, are distinct but not mutually exclusive roles that require specific obligations that enhance accountability, compliance, and certainty.

Considerations on AI’s Role in Decisionmaking

A core unresolved issue in determining the correct scope of AI regulation is balancing the breadth desired by advocates, who argue that regulation must cover all potential ways systems can lead to discriminatory outcomes, with the industry’s need for operational clarity. **So far, no approach has satisfied both groups or adopted language they can agree on.**

Industry representatives argue that broader thresholds, such as “facilitating decision-making,” could unintentionally regulate low-risk and essential technologies like calculators or Excel spreadsheets, which are not typically considered AI. Additionally, software that may “facilitate” a decision, such as scheduling tools, could also be affected, despite not having a legal or similarly significant impact on individuals.

In contrast, civil society and civil rights groups argue that narrow thresholds like “controlling factor” may enable organizations to evade regulatory responsibility by merely having humans rubber-stamp decisions, failing to meaningfully address the many ways in which AI systems can result in discriminatory outcomes.⁸

Even the “**substantial factor**” threshold, intended as a middle ground, is not supported by either group, with industry concerned about clarity and civil society worried about insufficient protections.

⁸ For example, an investigation of 391 employers’ compliance under NYC Local Law 144 (which uses “substantially assist”) found only 18 posted audit reports and 13 transparency notices. Lucas Wright et al., [“Null Compliance: NYC Local Law 144 and the Challenges of Algorithm Accountability,”](#) FAcT ‘24 at 1,701–13 (2024); Matt Scherer, [“Regulating Robo-Bosses: Surveying the Civil Rights Policy Landscape for Automated Employment Decision Systems,”](#) Center for Democracy and Technology (July 2024).

⁹ Though Connecticut was the first state to debut this specific language, [New York City Local Law 144](#) used “**substantially assist** or replace discretionary decision making.” Federally, [the Office of Management of Budget](#) defined “rights-impacting AI” as AI whose output serves as a **principal basis** for a decision or action.

Alternative Approaches: Regulating Government Entities

Most states that address private sector regulation typically start by focusing on the use of AI or automated systems by government agencies. When it comes to AI in "consequential decisions" and areas protected by U.S. civil rights law, government use of AI often includes critical areas such as access to government benefits and criminal justice. This approach is exemplified by [Maryland SB 818](#) (2024), [Connecticut SB 1103](#) (2023), and [Virginia SB 487](#) (2024).

Developer: Persons or entities that are developing, or creating, an AI system. When creating an AI tool or system, developers generally determine the intended purpose and scope of the AI system, gather and preprocess data to train the model, choose or design the appropriate algorithm or model architecture, train the model, and then conduct the necessary evaluation and optimization. If/when the AI system is sold or shared with a deployer, the developer may assist with integrating the application or system for real-world use, but often do not have ongoing access to deployer environments. Therefore, regulations typically focus developer obligations on testing systems they develop, providing necessary documentation about the system, and assisting deployers in their obligations.

Deployer: Persons or entities that are using AI systems for certain in-scope areas (such as employment, healthcare, or financing). Typically, deployers directly interact with individuals

and ultimately decide the context in how an AI system is used, and therefore deployer obligations in regulations generally focus on providing notice to affected individuals and conducting post-deployment monitoring. Deployers are also often required to maintain a risk management program and conduct their oversight and testing (discussed below).

Considerations on the Developer-Deployer Distinction:

Tailoring obligations to clearly defined roles is important for effective AI regulation, as the responsibilities of developers and deployers differ significantly. Developers, who design AI tools, have extensive control and visibility during the design phase. In contrast, deployers, who use AI tools in practice, have greater insight into the implementation and real-world use of these systems. For instance, a business deploying a commercially available AI tool may not have access to information about how the AI system was trained or designed unless the developer provides it. Conversely, developers' limited access to deployer environments—due to legal, contractual, privacy, and security constraints—contributes to the challenge of monitoring their products in practice, making it difficult to detect issues like algorithmic discrimination unless informed by deployers who conduct real-time testing.

Additionally, in some scenarios, an entity may be both a developer and deployer, subject to both role-specific obligations. Under many proposals, like the [Colorado AI Act](#), **a deployer can become a developer** if they engage in certain activities, such as making a substantial modification to an AI system. Open source models also raise particular challenges since their use involves numerous, sometimes unidentifiable parties, or entities that serve in both the developer and deployer role.

(5) Common Exceptions: Lawmakers drafting legislation to regulate high-risk AI systems have publicly considered a number of exceptions from applicability based on the nature of the technology at issue, the capacity of businesses, and interaction with existing laws. The four most common categories of exceptions are described below:

- **Technology Exceptions:** Many proposals explicitly exclude certain technologies from their definition of covered AI systems in order to avoid placing unnecessary requirements on the use of certain commonplace technologies, such as those that are inherently low risk and/or do not pose targeted risks to individuals. Examples of technologies and tools commonly excluded from coverage include calculators, databases and data storage, map navigation, spam and robo-call filtering, spellchecking, spreadsheets, and web hosting. Other categories of tools, such as cybersecurity and anti-fraud related technologies, may also be excluded from scope, as in the [Colorado AI Act](#) and [California AB 2930](#).
- **Existing Law Exceptions:** As both state and federal lawmakers develop AI regulations, proposals often exclude organizations and technologies already governed by existing sectoral laws to prevent overlapping compliance obligations. For instance, [Colorado's 2021 legislation](#) on predictive models in insurance ratings led the [Colorado AI Act](#) to exclude insurers and AI developers compliant with this existing statute. The [Colorado AI Act](#) also does not apply to systems that have been “approved, authorized, certified, cleared, developed, or granted by a federal agency” or that are “in compliance with standards established by a federal agency.”
- **Small Business Exceptions:** Lawmakers considering the regulation of high-risk AI systems have sought to avoid placing disproportionate regulatory burdens on small businesses and startups.¹⁰ Typically, these exclusions have focused on small businesses who serve as *deployers* of ‘off the shelf’ AI systems and who use that system for its intended purpose and do not make any significant updates to the system or train the system with their own data. Conversely, some lawmakers argue against small business exemptions, noting that, like existing civil rights laws, which do not exempt small businesses, they should also be held accountable for addressing algorithmic discrimination.
- **Public Interest Exceptions:** Finally, lawmakers have considered ‘public interest’ exceptions, such as those in the [Colorado AI Act](#), which exempts activities like taking immediate steps to protect an individual’s life or physical safety, engaging in research in the public interest, effectuating a product recall, identifying and repairing technical errors and conducting pre-deployment research, testing, and development activities of an AI system.

¹⁰ “[Artificial Intelligence Regulations April 18th Forum with State Legislators](#),” *supra* footnote 4 (noting that governance frameworks should balance incentivizing innovation and supporting local businesses and workforce).

B. Provisions Regarding Algorithmic Discrimination

The occurrence of algorithmic bias is well-documented and can pose a serious risk to the livelihoods of individuals.¹¹ Lawmakers who have adopted the ‘consequential decision’ approach to legislation have indicated that a primary goal is to mitigate discrimination that may arise from the use and application of AI against protected classes.¹² Therefore, most leading proposals have included specific provisions regarding algorithmic discrimination.

Definition of Algorithmic Discrimination: Though there are small differences among leading proposals, most laws and proposed laws reflect a consensus that algorithmic discrimination is defined as a condition where the use of an artificial intelligence system results in unlawful or unjustified differential treatment or impact that disfavors an individual or group of individuals on the basis of their actual or perceived protected class.¹³ It is also common for proposals to utilize exemptions set forth in the Lawyers Committee Model Online Civil Rights Act, which allows for certain activities, including self-testing for bias, activities that support increased diversification, and acts conducted by a private club that are currently exempted under civil rights law.¹⁴

Specific Provisions for Algorithmic Discrimination: Two approaches have emerged as leading models for legislating on algorithmic discrimination: one, utilized by [California AB 2930](#), which sets forth a blanket prohibition against algorithmic discrimination, and the other, utilized by the [Colorado AI Act](#) and [Connecticut SB 2](#), that creates a duty of care to prevent algorithmic discrimination.

- **Prohibition on Algorithmic Discrimination:** [California AB 2930](#) would prohibit deployers from using an automated decision tool and prohibit developers from making available an

¹¹ Bias as a mathematical phenomenon is inevitable. See Sigal Samuel, [Why It's So Damn Hard to Make AI Fair and Unbiased](#), Vox (April 19, 2022) (illustrating how different conceptions of bias and fairness can lead to different outcomes). The trick is to mitigate against harmful bias to ensure fairness. Nicol Turner Lee, Paul Resnick & Genie Barton, [Algorithmic bias detection and mitigation: Best practices and policies to reduce consumer harms](#), Brookings (May 22, 2019).

¹² “[Senator Maroney Leads Advancement of Landmark Artificial Intelligence Legislation, Encouraging Use and Developing Guardrails for Adoption](#),” Connecticut Senate Democrats (April 24, 2024). Individual harms that may arise from automated decision-making include loss of opportunity, economic loss, or loss of liberty. Future of Privacy Forum, [Unfairness by Algorithm: Distilling the Harms of Automated Decision-Making](#) (Dec. 2017).

¹³ Protected classes vary based on a state’s specific anti-discrimination laws and may include things like race, color, ethnicity, age, sex, disability, genetics, limited English proficiency, veteran status, and more.

¹⁴ The [Lawyers Committee Online Civil Rights Model Bill](#) (Dec. 2023) definition of “discrimination” exempts: (1) the offer, license, or use of a covered algorithm for the sole purpose of— (A) a developer’s or deployer’s self-testing to identify, prevent, or mitigate discrimination or otherwise to ensure compliance with obligations under federal law; or (B) expanding an applicant, participant, or customer pool to increase diversity or redress historic discrimination. (2) any private club or group not open to the public, as described in section 201(e) of the Civil Rights Act of 1964 (42 U.S.C. 2000a(e)).

automated decision tool if an impact assessment “identifies a reasonable risk of algorithmic discrimination.” Though earlier versions of [California AB 2930](#) (and its 2023 predecessor [AB 331](#)) that were relied upon in other state proposals broadly prohibited employers from using automated decision tools that resulted in algorithmic discrimination.

- **Duty of Care:** Under the [Colorado AI Act](#), both developers and deployers are subject to a duty to use “reasonable care” to protect consumers from “any known or reasonably foreseeable risks of algorithmic discrimination from the intended and contracted uses” of the high-risk AI system. Under this framework, developers and deployers maintain a rebuttable presumption of using reasonable care under this provision if they satisfy the obligations of the relevant statute.

Considerations on Algorithmic Discrimination:

- **Prohibition versus Duty of Care:** In practice, a blanket prohibition against algorithmic discrimination may be more likely to impose strict liability than a duty of care, which may be assessed using a proportionality test that considers factors, circumstances, and industry standards to determine whether an entity exercised reasonable care to prevent algorithmic discrimination. Conceptually, industry representatives have supported this approach, but consumer advocates have expressed concern that it departs from civil rights standards.¹⁵
- **Interaction with Existing Civil Rights Law:** Although most agree that civil rights laws already apply to AI systems in theory, civil rights experts note that the law is far behind the technology, leaving open questions regarding regulatory liability and guidance. Therefore, some civil society organizations and industry representatives agree that AI-specific laws with tailored algorithmic discrimination provisions could add needed clarity. However, it is unclear how these provisions will interact with existing civil rights law, particularly as it relates to disparate impact analysis, where seemingly neutral practices disproportionately affect one group of people with a protected characteristic more than another. Under existing law, a finding of disparate impact is not per se unlawful, but requires further analysis to determine if a practice is discriminatory. Federal regulators, data scientists, and civil rights advocates argue that disparate impact is a necessary component of ensuring AI non-discrimination.¹⁶

¹⁵ See, e.g., [“Civil Rights Standards for 21st Century Employment Selection Procedures,”](#) Center for Democracy and Technology, American Association for People with Disabilities, American Civil Liberties Union, The Leadership Conference on Civil and Human Rights, National Women’s Law Center, Upturn, (December 2022).

¹⁶ See, e.g., Olga Mack, [“Promoting AI Fairness: The Application of Disparate Impact Theory,”](#) MIT Computational Law Report (August 2023); [“Reflecting on Civil Rights and Our AI Future,”](#) The Leadership Conference on Civil and Human Rights (April 2023); [“Select Issues: Assessing Adverse Impact in Software, Algorithms, and Artificial Intelligence Used in Employment Selection Procedures Under Title VII of the Civil Rights Act of 1964,”](#) U.S. Equal Employment Opportunity Commission (May 2023).

C. Common Obligations for Developers and Deployers

Common developer and deployer obligations observed in state AI legislative and regulatory proposals include (1) **Transparency and Disclosures**; (2) **Assessments**; and (3) **AI Governance Programs**. See [Report Supplement](#), Table 3 for specific examples of language used in proposals.

Transparency: AI transparency refers to the practice of making the workings, decision-making processes, and impacts of AI systems clear and understandable to various stakeholders, including individuals and the public.

- **Notice To Individuals:** Notice to individuals typically requires that deployers provide certain information to people subject to automated decision-making technology or individuals interacting with AI products, often in the form of pre-use or just-in-time notice. These notices typically include information regarding what the AI system is used for, a general description of how the system works, and how an individual can exercise their relevant rights.
- **Notice to Public:** The ‘consequential decision’ approach also often includes public transparency requirements of high-risk AI systems in use by a deployer or developed by developer.
- **Documentation Between Developer and Deployer:** To further foster transparency in the AI lifecycle and ensure that deployers have the information needed to fulfill their obligations to consumers and the public, some proposals require certain documentation to be disclosed between developers and deployers.

Assessments: Risk or impact assessments and audits are tools that may be used to examine AI systems in different ways, including their performance, bias, and risk of discrimination. The use of assessments and audits are critical tools that can enhance trust, accountability, and the responsible deployment of AI technologies across various sectors and applications.¹⁷ While the terminology may lead some to think that these tools are singular there are substantive differences:

- **Risk/Impact Assessment:** Assesses and documents whether and to what extent an AI system poses a risk of discrimination to individuals. Like data privacy impact assessments, AI impact assessments typically include information regarding the system’s purpose, discrimination risks, benefits, and safeguards. Conducted properly, they should provide information that allows a developer or deployer to make a decision about the risk profile of an AI system and act accordingly with the law.

¹⁷ See, Dileep Srihari & Meghan Chilappa, [Impact Assessments: Supporting AI Accountability & Trust](#), Access Partnership, Workday (Jan. 28, 2023).

- **Audit:** Tests an AI system to evaluate technical aspects based on particular metrics such as accuracy and reliability, usually for bias based on protected characteristics. Audits alone are usually insufficient to make a decision about the risk profile of an AI system and usually need to be accompanied by an assessment process.

AI Governance Programs: Another core element of a comprehensive AI regulation is the creation of AI governance programs or risk management policies and procedures. These programs typically create a structured framework of policies, procedures, and controls designed to oversee and manage the development, deployment, and use of AI within an organization aimed at ensuring that AI technologies are developed and deployed responsibly, ethically, and in compliance with relevant laws and regulations. Common requirements of an AI governance program include:

- **Specifications:** It specifies the personnel, technical safeguards, and processes used to identify, document, and mitigate risks;
- **Reasonability:** It must be “reasonable” considering a variety of factors, including adherence to recognized frameworks (such as the National Institute of Standards and Technology AI Risk Management Framework),²¹ the size and complexity of the covered entity, the nature and scope of the system, and the volume and sensitivity of data processed in the system;
- **Iterative:** Programs should be planned, implemented, reviewed, and updated over the AI system’s lifecycle; and
- **Oversight:** At least one person or employee should be responsible for overseeing and maintaining the governance program.

Considerations on Audits:

Auditing AI systems can be important for identifying and mitigating both technical risks, such as accuracy and reliability, and socio-technical risks, such as bias and discrimination. However, the current lack of standardized methods and clear guidelines presents significant operational and policy challenges.¹⁸ This lack of standards can lead to inconsistent or ineffective audit methodologies, as seen with [New York City Local Law 144](#), which has been criticized for requiring potentially unreliable statistical measurements that may diverge from established standards like the ‘four-fifths’ disparate impact rule in existing civil rights law.¹⁹

While the National Institute of Standards and Technology is currently developing scalable auditing techniques,²⁰ lawmakers should also consider the need to build institutional capacity in both government and industry to manage and oversee the growing audit industry that may emerge from these requirements.

¹⁸ See Evan Selinger, Brenda Leong & Albert Fox Cahn, “[AI Audits: Who, When, How... Or Even If?](#)”, in COLLABORATIVE INTELLIGENCE: HOW HUMANS AND AI ARE TRANSFORMING OUR WORLD (forthcoming MIT Press 2024).

¹⁹ *Id.*

²⁰ “[NIST’s Responsibilities Under the October 30, 2023 Executive Order](#),” National Institute of Standards and Technology.

²¹ “[AI Risk Management Framework \(AI RMF 1.0\)](#),” National Institute of Standards and Technology (Jan. 2023).

D. Common Consumer Rights

AI ‘consequential decision’ frameworks often establish rights and protections for individuals interacting with or affected by AI systems, creating corresponding obligations for deployers directly engaging with these individuals. Lawmakers typically evaluate a range of rights to ensure that individuals have both awareness of and recourse against potential harms caused by AI systems. Beyond general transparency requirements that provide individuals with necessary information (detailed above), these frameworks may grant specific rights, including the **right to notice and explanation** about the use of an AI system, **the right to correct** inaccurate information used in decision-making, and the **right to appeal or opt-out** of an automated decision. Sample language for these rights is provided in [Report Supplement](#), Table 4.

Right to Notice and Explanation: Most ‘consequential decision’ approaches acknowledge the need for individuals to know that an automated system is being used before it is used on them, how and why it is being used, and/or when an adverse decision was rendered by the automated system. Laws written in these frameworks generally require that notices provided before an automated decision is used should be in plain language and accessible to individuals. Often, notice provisions require the following information to be disclosed to individuals:

- A statement that an automated system is used for decision-making in the specified context;
- A general description of the system’s purpose and nature of the decision;
- The deployer’s contact information;
- Instructions on how an individual can exercise their relevant rights and sufficient information to assist an individual in doing so.²²

Some proposals, such as [California AB 2930](#), require disclosure of additional details such as “(i) The personal characteristics or attributes that the automated decision tool will measure or assess[,] (ii) The method by which the automated decision tool measures or assesses those attributes or characteristics[,] (iii) How those attributes or characteristics contribute to the consequential decision . . . [and] (vi) A summary of the most recent impact assessment performed on the automated decision tool.”

If an adverse decision is made by an AI system, some existing sector-specific laws and the [Colorado AI Act](#), require additional information to be disclosed to the individual that explains the decision.²³ This transparency is important particularly if an adverse decision was incorrect and/or

²² For example, the [Colorado AI Act](#) obligates deployers who are also “controllers” under the Colorado Privacy Act to provide information about consumers’ data privacy rights in this notice. § 6-1-1703(4)(a)(III).

²³ For example, the Fair Credit Reporting Act requires anyone who uses a credit report to take an adverse action (such as denying an application for credit, insurance, or employment) to provide notice to the individual and provide information about the credit reporting agency. 15 U.S.C. § 1681m.

Application and Updates to Data Privacy Laws:

Instead of introducing stand-alone AI proposals, some lawmakers have included similar protections in newer state privacy laws. For example, the [Minnesota Consumer Data Privacy Act \(MNCDDPA\)](#) grants individuals the right to "question the result" of significant profiling decisions, akin to the "right to appeal" in the AI 'consequential decision' approach. It allows individuals to challenge profiling results, understand the reasons behind decisions, and learn about possible actions to achieve different outcomes.

Meanwhile, the [California Privacy Protection Agency \(CPPA\)](#) is in a [pre-rulemaking process](#) that would extend the application of the California Consumer Privacy Act to "automated decisionmaking technologies."

based on incomplete or inaccurate information in order for an individual to appeal that decision or exercise their right to correct the underlying information for the correct decision to be rendered. These adverse decision notices should be clear and accurate, and may include more specific information that helps the individuals understand how the specific decision was made

based on their particular information or circumstance. This additional information typically involves (i) **the principal reason** for the adverse decision; (ii) the degree in which the **automated system contributed** to the decision; (iii) the type of **data processed** to make the decision; and (iv) the **source of which the data** used for the decision was pulled from.

Right of Correction: Automated decision-making systems can make errors or rely on incorrect or outdated data that may unfairly impact individuals and their rights. Providing individuals with the right to correct this data is intended to help ensure decisions are based on accurate information, promoting fairness and preventing harm. With its origins in the rights to correct in data protection and privacy law and policy, AI 'consequential decision' approaches sometimes require deployers to provide individuals an opportunity to

correct any incorrect personal data processed in an automated decision system, if an adverse decision was rendered.²⁴

Right to Opt-out of Automated Decisions or Appeal: Across the state landscape, two core individual rights are emerging: opt-out rights not to be subject to automated systems and ex-post rights to appeal the results of a consequential decision produced by an AI system for human review. Notably, most state legislative frameworks opt for either establishing a right to appeal the result of an AI system or creating a right to opt out of being subject to a decision made by the system.

- **Opt-Out:** Proposals with opt-out rights allow individuals to request not to be subject to a covered decision made by an automated system before the system is deployed against them. The right is typically tied to a disclosure, **at or before a decision is made**, that provides an individual with meaningful information about the system. Where an individual

²⁴ See [Colorado SB 24-205](#) (enacted), § 6-1-1703(4)(b)(II).

opts out, the law may require that the individual receive an “alternative selection process or accommodation.”²⁵

- **Right to Appeal:** Proposals with a right to appeal allow individuals, after an adverse decision is reached, to request that the adverse result rendered by an automated system be reviewed by a human. The right is typically tied to a disclosure, **after a decision was made**, that details why a particular outcome was reached and what data was relied upon.

Considerations on Limited Exceptions to Carrying Out Individual Rights:

Sometimes, laws and proposals provide limited exceptions where a business is not required to comply with the request. Common exceptions include denying an appeal where compliance is not feasible or where honoring an appeal would not be “in the best interest” of a consumer, such as where a delay “might pose a risk to the life or safety” of the consumer.²⁶

E. Investigation and Enforcement by the State Attorney General

The majority of proposals and laws seeking to regulate high risk AI systems provide for exclusive enforcement by the state Attorney General’s office. Some proposals would allow other public enforcement bodies to also enforce the law, such as the state Civil Rights Department under [California AB 2930](#), or the state Commission of Consumer Protection under early versions of [Connecticut Senate Bill 2](#). Given areas of potential overlap with existing civil rights and consumer protection law, most frameworks aim to avoid preempting existing private rights of action and theories of liability, and some proposals explicitly state that while nothing in the Act is intended to give rise to an independent cause of action, existing causes of action are also preserved.²⁷ State AI legislative proposals have explored various mechanisms to enhance regulatory oversight and incentivize compliance with high-risk AI systems. These mechanisms include:

Regulatory Tools:

1. **Affirmative Reporting:** The [Colorado AI Act](#) mandates developers to disclose any “known or reasonably foreseeable risks of algorithmic discrimination” to the Attorney General and all known deployers or other developers of the high-risk AI system.
2. **Document Production:** State proposals often require developers and deployers of high-risk AI systems to maintain and produce documents such as risk management

²⁵ [California AB 2930](#) (proposed) (July 3, 2024), § 22756.2, subd. (b)(1); [New York City LL 144](#) (enacted), New York City, 2021, § 20-871(b)(1).

²⁶ See [California AB 2930](#) (proposed) (July 3, 2024), § 22756.2, subd. (b)(1); [Colorado SB 24-205](#) (enacted) (2024) § 6-1-1703(4)(b)(III).

²⁷ The [Colorado AI Act](#) and [Connecticut SB 2](#) both provide that nothing under the law “preempts or otherwise affects any right, claim, remedy, presumption, or defense available at law or in equity.” [Colorado SB 24-205](#) (enacted) (2024), § 6-1-1706(5); accord [Connecticut SB 2](#) (proposed) (Apr. 24, 2024), § 7(f)(4).

Alternative / Additional Approaches: Rulemaking

Lawmakers may delegate rulemaking authority to government regulators, like state Attorneys General, to establish more detailed, technical standards and allow for greater expert input and public consultation. For instance, under the [Colorado AI Act](#), the Attorney General has rulemaking authority over specific key requirements, including impact assessments and risk management frameworks. [Illinois HB 3773](#), however, broadly delegates rulemaking authority to the Department of Human Rights for the “implementation and enforcement” of the law.

Alternatively, state agencies may use rulemaking to adapt existing laws to AI-specific challenges. Rather than drafting new laws, regulators may look to update privacy laws (as seen with [CPPA automated decisionmaking regulations](#)) or civil rights laws (as seen with [California Civil Rights Council's proposed modifications](#) to employment regulations on automated decision systems).

policies and impact assessments. Organizations may be required to submit this documentation to enforcement agencies upon request, typically within 30 to 90 days.

Enforcement Mechanisms:

1. **Right to Cure:** [California AB 2930](#) allows regulators to provide written notice of alleged violations and offers a 45-day period for entities to address the issues. If the violation is cured within this timeframe and the entity provides a sworn statement affirming the cure, the regulator cannot pursue injunctive relief.

2. **Rebuttable Presumption:** Frameworks like those in [Connecticut](#) and [Colorado](#) offer a “rebuttable presumption” that a business has met its obligation to protect individuals if it completes required documentation and notifications.

3. **Affirmative Defense:** [Connecticut](#) and [Colorado](#) frameworks establish an affirmative defense for developers who (1) discover a violation, (2) cure it within 30 days and provide notice and evidence of the cure, and (3) comply with recognized risk management frameworks, such as those by the [National Institute of Standards and Technology](#) or the [International Organization for Standardization](#).

Such requirements typically provide certain exceptions for trade secrets and exclude such records from production under open records laws.²⁸ These requirements also generally provide that the required disclosures do not constitute a waiver of any attorney-client or work-product privilege or protection.²⁹

Considerations on Enforcement:

Most state lawmakers have been hesitant to include a private right of action in AI and data privacy bills. Consumer rights and civil rights organizations argue that such a right is crucial for holding technology accountable, especially concerning civil rights. On the other hand, industry advocates fear that a private right of action could lead to excessive litigation and place undue burdens on businesses and the legal

²⁸ See [Colorado SB 24-205](#) (enacted), §§ 6-1-702(7) and 6-1-1703(9); [California AB 2930](#) (proposed) (July 3, 2024), § 22756.8, subd. (b)(3)(B).

²⁹ See [Colorado SB 24-205](#) (enacted), §§ 6-1-702(7) and 6-1-1703(9); [California AB 2930](#) (proposed) (July 3, 2024), § 22756.8, subd. (b)(2).

system. Despite these concerns, there are varying degrees of private right of action provisions, including options that include rights to cure, require proof of injury, or operate pursuant to a sunset period.

III. Alternative Technology-Specific Approaches

While the most common approach to AI regulation in U.S. states—covering both government use and private sector—thus far has been a general risk-based approach, some lawmakers have pursued technology-specific regulations aimed at addressing risks of harm arising from specific forms of AI. Lawmakers have most commonly focused on **Generative AI Systems** (AI that can create new content such as text, images, music, or videos)³⁰ and **Frontier AI or Foundation Models** (large AI models that can be used in a wide variety of use cases and applications, sometimes referred to as “general-purpose AI”).³¹

When proposing regulations for generative AI, lawmakers have primarily focused on enhancing transparency about its use and outputs. Legislation often includes requirements for disclosure, such as providing consumer notices about the use or creation of generative AI, labeling content as synthetic or AI-generated, or implementing watermarking to clearly identify generated content.³² In 2024, Utah enacted [SB 149](#), requiring individuals or entities to clearly and conspicuously disclose when a generative AI system is interacting with a consumer, rather than a human. In 2024, the California legislature passed two bills aimed at increasing transparency for generative AI systems, pending the governor's signature. [California SB 942](#) would require entities providing generative AI tools to offer an “AI detection tool” that lets individuals check whether content was created or modified by the AI system. [California AB 2013](#) would mandate that developers of generative AI systems publicly disclose documentation about the data used to train these systems. California also has a [2018 law](#) that prohibits using a “bot” to communicate or interact with the intent to mislead individuals about the bot’s artificial identity.

A growing number of technology-specific proposals are now targeting Frontier AI or Foundation models. These models, due to their significant scale and power, can serve as the foundation for a wide range of AI applications. Consequently, some lawmakers have sought to regulate these models to ensure they are developed and deployed with robust safety measures and protocols, preventing misuse and mitigating unintended consequences. A key bill, [California SB 1047](#), passed by the legislature in 2024 and pending the governor's signature as of publication of this report, would require developers to certify to the government that their frontier models have a

³⁰ See Cole Stryker & Mark Scapicchio, “[What Is Generative AI?](#)”, IBM (Mar. 22, 2024).

³¹ See Elliot Jones, “[What Is a Foundation Model?](#)”, Ada Lovelace Institute (July 17, 2023).

³² See e.g., [California SB 942](#) (proposed) (Aug. 19, 2024); [California AB 3211](#) (proposed) (June 24, 2024) (relating to provenance, authenticity and watermarking standards for generative AI systems); [Massachusetts HB 4788](#) (proposed) (Jan. 22, 2024); and [Ohio SB 217](#) (proposed) (Jan. 24, 2024).

written “safety and security protocol” and the capability to promptly enact a “full shutdown” if needed.

Considerations on Frontier/Foundation Models:

Lawmakers in the U.S. have faced challenges in attempting to regulate frontier or foundation models, as evidenced by the controversy surrounding California SB 1047. These challenges stem from the complexity and scale of these models, their diverse range of applications, and the technical expertise required to develop effective regulatory standards. Additionally, derivative AI systems can be built based on foundation models through a process of “fine-tuning” the model to perform specific functions. Two major questions have arisen for policymakers: which models should be subject to new regulations, and how to ensure that new rules do not stifle the open-source ecosystem.

- **Computing Power Thresholds:** Because some experts argue that high complexity models may inherently pose greater risks,³³ various regulatory regimes have sought to place default restrictions on models trained using a certain amount of computing power. The [European Union AI Act](#) defines general purpose models to include those trained on computing power of more than 10^{25} floating-point operations (FLOPs). The [Biden Administration Executive Order on AI](#) places certain reporting requirements on models trained using computing power greater than 10^{26} FLOPs, a threshold that was replicated in California SB 1047. However, critics argue that computational power alone is not a reliable indicator of risk.³⁴ They contend that such approaches fail to address many high-compute systems currently in use and focus on speculative risks rather than evidenced risks, such as algorithmic bias.³⁵
- **Supporting Open Source:** Placing requirements and liability on the developers of foundation models may limit the ability for developers to make open-source models generally available for use and modification. The National Telecommunications Information Administration recently observed that “‘Open-weight’ models allow developers to build upon and adapt previous work, broadening AI tools’ availability to small companies, researchers, nonprofits, and individuals.”³⁶ Responding to concerns about the impact on open source, California SB 1047 received late amendments to exclude artificial intelligence models created by fine-tuning a covered model using computing power that costs less than ten million dollars.

IV. Conclusion

Advancements in automated systems promise groundbreaking insights and increased economic efficiency, but they also introduce significant challenges that demand intricate legal, technical, and social solutions. As artificial intelligence plays an increasingly central role in our lives—particularly in decisions affecting access to essential opportunities—lawmakers across the U.S. are actively pursuing strategies to harness AI’s benefits while addressing risks such as

³³ See Markus Anderljung et al., “[Frontier AI Regulation: Managing Emerging Risks to Public Safety](#)” (Nov. 7, 2023) (describing the various ways in which frontier AI models “pose a distinct regulatory challenge”).

³⁴ Ingrid Stevens, “[Regulating AI: The Limits of FLOPs as a Metric](#),” Medium (May 1, 2024).

³⁵ Cf. Alex Hanna & Emily M. Bender, “[AI Causes Real Harm. Let’s Focus on That over the End-of-Humanity Hype](#),” Scientific American (Aug. 12, 2023).

³⁶ “[NTIA Supports Open Models to Promote AI Innovation](#),” National Telecommunications and Information Administration (July 30, 2024).

algorithmic discrimination, privacy, and transparency. The complexity of the current legislative landscape underscores the nuanced challenges in advancing AI regulations. However, the emerging trends reveal a collaborative push for an interoperable framework, where consistent definitions and principles will be crucial in supporting business compliance, safeguarding individual rights, and ensuring regulatory clarity.

Special Acknowledgments:

We extend our sincere gratitude to the individuals who reviewed and provided invaluable feedback on this report. Their insights have greatly contributed to ensuring that the report reflects a broad spectrum of perspectives and considerations. The contributions of these individuals, representing a range of viewpoints, have been instrumental in shaping a more comprehensive and balanced analysis.

Note that their names do not imply endorsement of this Report. We appreciate their engagement and commitment to advancing state AI policy.

Connecticut State Senator James Maroney

Evangelos Razis and Lev Sugarman, Workday

Grace Gedye, Consumer Reports

Matt Scherer, Center for Democracy and Technology

Meghan Pensyl, BSA The Software Alliance

*If you have any questions, please contact **Tatiana Rice** (trice@fpf.org), **Jordan Francis** (jfrancis@fpf.org), or **Keir Lamont** (klamont@fpf.org).*

Disclaimer: This report is for informational purposes only and should not be used as legal advice.



1350 Eye Street NW Suite 350
Washington, DC 20005

info@fpf.org

FPF.org