

IN THE UNITED STATES DISTRICT COURT
FOR THE DISTRICT OF COLORADO

Civil Action No.: 1:26-cv-01515

X.AI LLC,

Plaintiff,

v.

PHILIP J. WEISER, COLORADO ATTORNEY GENERAL,

Defendant.

COMPLAINT FOR DECLARATORY AND INJUNCTIVE RELIEF

Plaintiff X.AI LLC (“xAI”) brings this civil action against Philip. J. Weiser, in his official capacity as Colorado Attorney General, for declaratory and injunctive relief and alleges as follows:

INTRODUCTION

1. Technologies powered by artificial intelligence (“AI”) are fast becoming the primary instrument through which we discover, verify, and transmit knowledge. AI systems are promising instruments of enlightenment and free expression, capable of democratizing knowledge in ways the printing press or the internet never could and accelerating scientific breakthroughs at a pace once unimaginable. They also have the potential to undermine free inquiry and public discourse on an unprecedented scale: whoever owns or aligns the dominant models can quietly alter what counts as “truth”—portraying ideological fads or falsehoods as reality—can filter inconvenient facts, and can curate entire worldviews. This is dangerous to the advancement of

human knowledge. Only by seeking objective truth can AI faithfully help humanity understand the universe and advance the progress of human civilization.

2. xAI is a leading frontier AI developer whose flagship model is known as Grok. xAI believes that AI's knowledge should be all-encompassing and as far-reaching as possible. It builds AI specifically to advance human comprehension and capabilities. xAI named Grok after a word in Robert Heinlein's famous work *Stranger in a Strange Land*. As Heinlein famously put it, to "grok" something is "to understand so thoroughly that the observer becomes a part of the observed—to merge, blend, intermarry, lose identity in group experience." Robert A. Heinlein, *Stranger in a Strange Land* 206 (Berkley Medallion Ed., 1968).

3. In keeping with that mission, xAI has designed and developed Grok to answer to only evidence and reason, without regard to political correctness, ideological biases, or anything that might distort objective truth. This unwavering commitment ensures that Grok discharges its fundamental mission—assisting humanity in understanding the universe. But the State of Colorado now seeks to force xAI to abandon its disinterested pursuit of truth and instead promote the State's ideological views on various matters, racial justice in particular.

4. This case seeks to enjoin the enforcement and implementation of Colorado's Senate Bill 24-205 ("SB24-205"), a statute that severely burdens the development and use of AI. SB24-205's purported focus is a prohibition on so-called "algorithmic discrimination." But SB24-205 is decidedly not an anti-discrimination law. It is instead an effort to embed the State's preferred views into the very fabric of AI systems. Its provisions prohibit developers of AI systems from producing speech that the State of Colorado dislikes, while compelling them to conform their speech to a State-enforced orthodoxy on controversial topics of great public concern. This attempted coercion

is unconstitutional under the First Amendment. Indeed, “[o]n the spectrum of dangers to free expression, there are few greater than allowing the government to change the speech of private actors in order to achieve its own conception of speech nirvana.” *Moody v. NetChoice, LLC*, 603 U.S. 707, 741-42 (2024).

5. SB24-205 has been controversial and legally suspect from the outset. When Governor Polis signed the law in May 2024, he expressed “reservations” about its onerous provisions.¹ Governor Polis cautioned that he was “concerned about the impact this law may have on an industry that is fueling critical technological advancements across our state for consumers and enterprises alike. Government regulation that is applied at the state level in a patchwork across the country can have the effect to hamper innovation and deter competition in an open market.”² Weeks later, Governor Polis, Attorney General Weiser, and Senator Robert Rodriguez (who introduced SB24-205) issued a joint statement announcing their desire to “revise the new law” and “minimize unintended consequences associated with its implementation.”³

6. A year passed and the General Assembly had not amended the law, prompting Governor Polis, Attorney General Weiser, U.S. Senator Bennet, U.S. Representatives Neguse and Pettersen, and Denver Mayor Johnston to ask the General Assembly in May 2025 to “delay implementation of SB 24-205 until January 2027” so that stakeholders could craft a solution that

¹ Jared Polis, Statement on Signing SB24-205 (May 17, 2024), <https://drive.google.com/file/d/1i2cA3IG93VViNbZxu9LPgbTrZGqhyRgM/view>.

² *Id.*

³ Jared Polis et al., Letter Concerning SB24-205 (June 13, 2024), https://drive.google.com/file/d/1UtIHZRKDq4a09q0mzT_o4cBj13_qanrg/view?usp=sharing.

does not “stifl[e] innovation or driv[e] business away from our state.”⁴ Attorney General Weiser repeated these concerns just three months later, in August 2025, acknowledging that the “bill is really problematic [and] needs to be fixed.”⁵ The Governor issued an executive order that same month convening a Special Session of the General Assembly in part to do just that.⁶ But in that session the legislature merely delayed SB24-205’s effective date to June 30, 2026—it did nothing to “fix” SB24-205. As of the date of this Complaint, the General Assembly has not introduced a bill to amend the law during the current legislative session, which ends the second week of May of this year.

7. SB24-205’s problems are many, but its most obvious constitutional defect appears in its first section. The law defines “algorithmic discrimination” as “any condition in which the use of an [AI] system results in an unlawful differential treatment or impact that disfavors an individual or group.” § 6-1-1701(1)(a).⁷ But the law then simultaneously *promotes* a form of “differential treatment” the State favors—discrimination intended to “increase diversity or redress historical discrimination.” § 6-1-1701(1)(b)(I)(B).

8. SB24-205’s internally contradictory, politicized definition of “algorithmic discrimination” underlies nearly all of its substantive requirements. Among other things, the law

⁴ Jared Polis et al., Letter to Colo. General Assembly (May 5, 2025), <https://drive.google.com/file/d/1YsQalh778UIkHKQNXwg1WMEtWqtUWYCR/view?usp=sharing>.

⁵ Cameron Marx, *State AG Warns Colorado AI Bill Could Drive Innovation Out of State*, BroadbandBreakfast: AI (Aug. 5, 2025), <https://broadbandbreakfast.com/state-ag-warns-colorado-ai-bill-could-drive-innovation-out-of-state/>.

⁶ Colo. Exec. Order No. D 2025 009 (Aug. 6, 2025), <https://drive.google.com/file/d/1kx5-WNwRYMD7K33lJQEqaSm6RsMUBb0X/view?pli=1>.

⁷ Unless otherwise stated, statutory citations refer to the Colorado Revised Statutes.

imposes affirmative duties on AI developers to “protect” consumers from “risks of algorithmic discrimination.” §§ 6-1-1702(1), 6-1-1703(1). It mandates disclosures about those “risks” to other “developers” and “deployers” of AI systems, to consumers, to the public, and to the Attorney General. §§ 6-1-1702(2)-(5), 6-1-1702(7), 6-1-1703(4)-(5), 6-1-1703(9). And the law mandates the implementation of risk-management policies and the creation of “impact assessments.” § 6-1-1703(2)-(3).

9. By requiring “developers” and “deployers” to differentiate between discrimination that Colorado disfavors and discrimination that Colorado favors, SB24-205 compels Plaintiff xAI—a “developer” under the law—to alter Grok, forcing Grok’s output on certain State-selected subjects to conform to a controversial, highly politicized viewpoint. But the State “may not compel [xAI] to speak its own preferred messages.” *303 Creative LLC v. Elenis*, 600 U.S. 570, 586 (2023). Doing so impermissibly “alter[s] the expressive content of [xAI’s speech].” *Id.* at 585 (quoting *Hurley v. Irish-Am. Gay, Lesbian and Bisexual Grp. of Boston*, 515 U.S. 557, 572-73 (1995)). And by “[m]andating speech that [xAI] would not otherwise make,” SB24-205 “necessarily alters the content of the speech.” *Riley v. Nat’l Fed. of the Blind*, 487 U.S. 781, 795 (1988). SB24-205’s content- and viewpoint-based speech compulsion is a clear violation of xAI’s First Amendment rights.

10. The First Amendment is not the only constitutional provision SB24-205 violates. The law broadly applies to “any use of an [AI] system to generate any content, decision, prediction, or recommendation concerning a consumer that is used as a basis to make a consequential decision”—a decision affecting education, employment, financial services, and more. § 6-1-1701(3), § 6-1-1704(11)(b). And it defines the “consumer” it purports to protect as “an[y]

individual who is a Colorado resident.” § 6-1-1701(4). In other words, SB24-205 is not geographically limited to Colorado; despite the fact that AI systems are deployed across the country, SB24-205 applies anywhere “a Colorado resident” is affected by an AI system. SB24-205’s provisions therefore apply to a wide range of AI uses nationally—regardless of where the “consumer” is located, where the AI system is developed, deployed, or used, or how foreseeable the AI use is—so long as a single “Colorado resident” is impacted by the AI system. As applied to xAI, a Nevada-incorporated, California-headquartered company with no offices in Colorado, SB24-205 violates the Dormant Commerce Clause by directly regulating development and deployment activities occurring entirely outside Colorado. It also fails the balancing test established by *Pike v. Bruce Church, Inc.*, 397 U.S. 137 (1970), by imposing burdens on interstate commerce that far outweigh any local benefit.

11. SB24-205 is also unconstitutionally vague. It fails to adequately define essential terms such as “high-risk artificial intelligence system,” “algorithmic discrimination,” and “historical discrimination.” The General Assembly left it to the Colorado Attorney General—who has exclusive authority to enforce SB24-205’s provisions—to decide what those provisions mean by “promulgat[ing] rules.” § 6-1-1707(1). This vagueness invites arbitrary enforcement.

12. Finally, SB24-205 codifies discrimination without legally sufficient justification. The bill—which lacks any legislative findings about the “algorithmic discrimination” the General Assembly elected to proscribe—compels AI developers like xAI to endorse Colorado’s views on diversity, equity, and inclusion or face significant compliance costs and civil fines. The State-prescribed orthodoxy SB24-205 requires is as harmful as it is unconstitutional. It forces AI

developers to distort their AI models to seek and output progressive ideology instead of the truth. SB24-205 therefore violates the Equal Protection Clause.

13. As Governor Polis and Attorney General Weiser have publicly (and repeatedly) acknowledged, SB24-205 is fundamentally flawed. Unless the implementation and enforcement of SB24-205 is enjoined, it will violate xAI's constitutional rights and cause irreparable constitutional harm, impose enormous burdens on xAI and the AI industry, and substitute Colorado's political preferences for the national economic and security imperative of American AI dominance.

14. The Court should declare SB24-205 unconstitutional and enjoin the Attorney General from enforcing it against xAI.

JURISDICTION & VENUE

15. This Court has subject-matter jurisdiction under 28 U.S.C. §§ 1331 and 1343(a). This Court has authority to grant legal and equitable relief under 42 U.S.C. § 1983, injunctive relief under 28 U.S.C. § 1651, and declaratory and other appropriate relief under 28 U.S.C. §§ 2201(a) and 2202.

16. Venue is proper in this District under 28 U.S.C. § 1391(b) because Attorney General Weiser resides in Colorado, and the events giving rise to this action occurred here.

PARTIES & STANDING

17. xAI is a limited liability company organized under the laws of Nevada with its principal place of business in Palo Alto, California. xAI maintains no offices in Colorado.

18. xAI develops AI models, and, in particular, large language models. xAI's signature product is a general-purpose large language model called Grok. Users rely on Grok for a vast array

of purposes, including drafting professional emails, analyzing business data, and answering questions using current information. But Grok’s potential uses extend far beyond these common applications. Human resources professionals can use Grok to screen and summarize resumes, generate interview questions, and analyze candidates’ social media profiles. Mortgage brokers can use Grok to process documents, run fraud checks, and expedite the underwriting process. Healthcare providers can use Grok to draft discharge summaries, synthesize medical literature, and assist in reviewing medical images.

19. Grok is a “high-risk artificial intelligence system” as defined by SB24-205, § 6-1-1701(9)(a), because individuals and entities use Grok to make, or rely on it as a substantial factor in making, “consequential decisions.” A “[c]onsequential decision” is one that “has a material legal or similarly significant effect on the provision or denial to any consumer of, or the cost or terms of,” numerous services or opportunities, including education, employment, finance, government services, healthcare, housing, insurance, and legal services. § 6-1-1701(3).

20. xAI is also a “developer” as defined by SB24-205, § 6-1-1701(7), because it does business in Colorado. xAI makes Grok available to individuals, entities, and governments nationwide through a variety of consumer- and enterprise-focused products. xAI’s customers include Colorado individuals and businesses. The consumer version of Grok, accessible through a standalone app or at Grok.com, has many registered users in Colorado.

21. If SB24-205 takes effect on June 30, 2026, xAI will have to comply with the law’s many onerous requirements for developers. These include a duty to exercise “reasonable care to protect consumers from any known or reasonably foreseeable risks of algorithmic discrimination.” § 6-1-1702(1). They also include extensive disclosure obligations. xAI must disclose to

deployers—entities doing business in Colorado that “deploy[] a high-risk [AI] system,” § 6-1-1701(6)—and to the public its practices for evaluating and mitigating “algorithmic discrimination.” § 6-1-1702(2)-(4). xAI must also disclose to the Attorney General, “in a form and manner” of his choosing, risks of “algorithmic discrimination” if such discrimination is “reasonably likely” to occur, has occurred, or if xAI receives a “credible report” that it has occurred. § 6-1-1702(5), (7).

22. SB24-205 imposes severe legal consequences for noncompliance. Violations “constitute[] an unfair trade practice,” § 6-1-1706(2), carrying a penalty of \$20,000 per violation, § 6-1-112(1)(a). The Attorney General may also pursue injunctive remedies and fee-shifting.

23. xAI has standing to bring this action. To allege standing, a plaintiff need only allege facts sufficient to show that its “conduct is ‘arguably ... proscribed by [the] statute’ [it] wish[es] to challenge.” *Susan B. Anthony List v. Driehaus*, 573 U.S. 149, 162 (2014) (quoting *Babbitt v. United Farm Workers Nat’l Union*, 442 U.S. 289, 298 (1979)); *see also id.* at 158-59 (“[W]e do not require a plaintiff to expose himself to liability before bringing suit.” (quoting *MedImmune, Inc. v. Genentech, Inc.*, 549 U.S. 118, 128-29 (2007))); *Virginia v. Am. Booksellers Ass’n*, 484 U.S. 383, 393 (1988) (finding standing where plaintiffs “alleged an actual and well-founded fear that the law will be enforced against them”).

24. xAI is a “developer” of a “high-risk artificial intelligence system” under SB24-205. The statute thus directly regulates xAI and xAI will face both substantive and disclosure obligations if the law takes effect. xAI has an “actual and well-founded fear that the law will be enforced against” it if it does not comply. *Am. Booksellers*, 484 U.S. at 393. “The State has not suggested that the newly enacted law will not be enforced,” and there is “no reason to assume

otherwise.” *Id.* Compliance with SB24-205 would violate xAI’s constitutional rights as set forth below. That is an injury-in-fact sufficient to confer standing. *See, e.g., Ward v. Utah*, 321 F.3d 1263, 1267 (10th Cir. 2003).

25. Defendant Philip J. Weiser is the Colorado Attorney General, sued in his official capacity. As the State’s chief law-enforcement officer, Attorney General Weiser is charged with enforcing Colorado’s laws, including SB24-205. The statute expressly provides that the Attorney General “has exclusive authority to enforce” its provisions, and that it “does not provide the basis for, and is not subject to, a private right of action.” § 6-1-1706(1), (6).

26. The Attorney General also has authority to implement SB24-205 by “promulgat[ing] rules as necessary for the purpose of implementing and enforcing” its provisions. § 6-1-1707(1).

27. Because Attorney General Weiser is statutorily directed to enforce SB24-205 and to promulgate rules implementing its provisions, and has not disclaimed any intent to do so, xAI has a credible fear that he will enforce the statute against the company. Attorney General Weiser is therefore a proper defendant against whom xAI may seek prospective relief to enjoin the implementation and enforcement of SB24-205. *Free Speech Coal., Inc. v. Anderson*, 119 F.4th 732, 735 (10th Cir. 2024).

BACKGROUND

A. Artificial Intelligence and Large Language Models

28. Computers are largely rule- and logic-based systems. Historically, programmers created algorithms that explicitly defined—*i.e.*, hard-coded—every step, scenario, and rule necessary to carry out the computer’s functions. Most AI systems take a radically different

approach: they use methods and techniques that enable computers to simulate human intelligence by capturing and learning from complex patterns and relationships in data to solve problems and achieve desired results—without requiring programmers to compile predefined logical rules that account for all possible scenarios.

29. Large language models (“LLMs”), like xAI’s flagship AI system Grok, are among the most powerful AI models available today. LLMs are trained on diverse datasets and designed to distill the patterns and structures of natural language. They generate text in response to “prompts”—questions supplied by users. An LLM’s goal is to produce a sensible response, referred to as an “output.” For example, if a user inputs the prompt “We saw the Golden Gate Bridge on our trip to,” the LLM may produce “San Francisco” as output.

30. To build an LLM, engineers first develop a computational model capable of learning patterns in language. These models—sometimes called “neural networks”—learn patterns and relationships between words and phrases in the training data. They do so through an iterative process, without being explicitly programmed with human-defined rules. What the models learn is reflected in a matrix of numerical values called “weights.” Weights map relationships among words and concepts, enabling the model to compute, for example, how strongly the concept “dog” is associated with “puppy” or “mutt” rather than “cat.” After training, the model uses these weights to analyze user inputs and generate outputs.

31. There are two phases to training an LLM: **pretraining** and **fine-tuning**. Pretraining is the initial phase in which the LLM is trained on text data to recognize general patterns in language and knowledge. The objective is to generalize across many types of tasks, contexts, and language structures.

32. Pretraining begins with assembling a training corpus. See Yiheng Liu et al., *Understanding LLMs: A Comprehensive Overview from Training to Inference*, 620 *Neurocomputing C*, at 6 (2024), <https://arxiv.org/pdf/2401.02038> (Jan. 6, 2024 preprint). AI models are trained on a wide array of datasets covering a vast range of information and data formats (e.g., text, images, and audio). Developers emphasize using training data that is diverse in subject matter and linguistic style because LLMs trained on narrow data perform poorly: they can handle only a limited range of tasks and have a more limited understanding of facts and language. See Brando Miranda et al., *Beyond Scale: The Diversity Coefficient as a Data Quality Metric for Variability in Natural Language Data*, arXiv preprint 1 (July 2, 2025) (“In particular, the richness and variety of data, otherwise known as data diversity, is a key aspect of data quality that plays an important role in researchers’ and practitioners’ choice of pre-training corpora for general capabilities. In other words, diverse data is high quality data when your goal is to instill general capabilities in a model.” (internal citations omitted)), <https://arxiv.org/pdf/2306.13840>.

33. Assembling the training corpus requires decisions about which domains and data sources to include or exclude. See Wayne Xin Zhao et al., *A Survey of Large Language Models*, arXiv preprint 17-18 (Mar. 18, 2026) (describing common sources of training data and highlighting the importance of removing low-quality data from the training corpus), <https://arxiv.org/pdf/2303.18223>. Identifying high-quality, diverse data facilitates learning, enables the corpus to capture the complexity and variety of human language, and allows the LLM to generalize more effectively.

34. During pretraining, the LLM undergoes a computation-intensive process in which it attempts to predict the next token (*i.e.*, a numerical representation that corresponds to a word,

part of a word, or particular character) in a sequence. See Yiheng Liu et al., *Understanding LLMs: A Comprehensive Overview from Training to Inference*, 620 Neurocomputing C, at 10 (2024) (explaining that during pretraining “the model is required to predict the next word in a given context” and that this “enables the model to develop a nuanced understanding of language”), <https://arxiv.org/pdf/2401.02038> (Jan. 6, 2024 preprint). At each step, the model takes a sequence of text from the training corpus and removes the last token. The LLM then predicts the removed token. As the neural network makes its prediction, each weight is tracked to compute its effect on the outcome. The prediction is compared to the actual deleted token. If correct, the weights are adjusted to reinforce the prediction; if incorrect, the weights are updated to minimize error. This process repeats until the LLM has made predictions for all tokens in the training corpus.

35. For example, suppose the sequence “We saw the Golden Gate Bridge on our trip to San Francisco” appears in the training corpus. During pretraining, all tokens except the final one (“We saw the Golden Gate Bridge on our trip to San”) are fed into the LLM. The model might incorrectly predict “Jose” as the next token. That prediction is compared to the correct token, “Francisco,” and the weights are adjusted to improve subsequent predictions.

36. This iterative process yields a multidimensional matrix of weights reflecting the information gleaned from the training corpus. These weights capture nuances in meaning, word relationships, grammatical properties, and essentially all characteristics of language necessary to approximate human language use. This matrix of weights, together with the code that runs the model, constitutes the LLM’s foundation model.

37. Although the foundation model is a functioning LLM, it may exhibit unintended behaviors, such as fabricating facts or ignoring user instructions. Output quality can be improved

through “fine-tuning,” a secondary training phase in which the model is refined for specific tasks, domains, or performance goals. See Haitao Jiang et al., *Supervised Fine-Tuning Versus Reinforcement Learning: A Study of Post-Training Methods for Large Language Models*, arXiv preprint 1 (Mar. 14, 2026) (explaining that “LLMs often require task-specific post-training adaptation to improve accuracy, mitigate erroneous outputs, and handle new tasks”), <https://arxiv.org/pdf/2603.13985>. Unlike pretraining, which builds foundational knowledge, fine-tuning adapts a pretrained model using smaller, high-quality datasets. Two common fine-tuning methods are supervised learning and reinforcement learning.

38. In supervised learning, the model is further trained on datasets of prompt-response pairs, where each prompt is paired with a high-quality response created or validated by humans or expert systems. See Bolin Zhang et al., *A Survey on Data Selection for LLM Instruction Tuning*, arXiv preprint 1 (Aug. 26, 2025) (explaining that supervised learning “aligns model outputs with human intent” using “curated (instruction, response) pairs”), <https://arxiv.org/pdf/2402.05123>. Constructing these datasets is an intensely editorial undertaking. Developers select prompts reflecting their views about what the model should learn. Some prioritize prompts testing legal reasoning or creative writing; others focus on prompts eliciting harmful responses so the model learns to refuse them. Developers must also decide what constitutes an “ideal” response-balancing competing values. Should the model respond to a complex legal prompt (e.g., “Was *Marbury v. Madison* rightly decided?”) with every relevant citation, or only key citations? Should it respond to a controversial prompt (e.g., “Does climate change increase wildfire frequency?”) with evidence-based findings alone, or with a narrative deemed politically correct? These judgment calls become embedded in the model’s weights.

39. Consider a simplified example: the relationship between “Jokić” and “Basketball” (a “vector”) might be assigned a weight of 0.6 in the foundation model, reflecting the probabilistic frequency of their association. The vector for “Gilgeous-Alexander” and “Basketball” might receive a lower weight, say 0.3. When prompted “Who is the best basketball player?” the model assesses billions of vectors and weights and outputs a response based on those frequencies—perhaps “Jokić is the best basketball player.” But an AI developer from Oklahoma City who dislikes the Denver Nuggets could train its AI model on curated datasets that favor Gilgeous-Alexander, such that the developer’s model subsequently attributes 0.9 to the weight for “Gilgeous-Alexander” and “Basketball” and causes the model to respond “Gilgeous-Alexander is the best basketball player.”

40. Developers also use reinforcement learning to assess performance, improve processing, and align the model toward preferred behavior. Through reinforcement learning, the AI system learns by trial and error: developers provide feedback as rewards (for good outputs) or penalties (for bad ones). *See* Long Ouyang et al., *Training Language Models to Follow Instructions with Human Feedback*, arXiv preprint 2, 18 (Mar. 4, 2022) (explaining that reinforcement learning uses a reward-function to improve model performance and describing editorial choices that inform the alignment process), <https://arxiv.org/pdf/2203.02155>. In response to a single input, the system may produce multiple responses, ranked by humans or by reward models trained on human feedback. Developers guide this process with rubrics detailing how annotators should evaluate and rank outputs. *See* Amelia Glaese et al., *Improving Alignment of Dialogue Agents via Targeted Human Judgements*, arXiv preprint 4-6 (Sept. 28, 2022) (describing one developer’s approach for writing rules to guide annotators in the reinforcement learning process),

<https://arxiv.org/pdf/2209.14375>. This process necessarily involves editorial decision-making. One developer might reward unvarnished candor on politically sensitive topics; another might penalize any output perceived as insensitive, regardless of empirical grounding.

41. Yet another way in which AI developers can input value judgments into an AI system is through system prompts. System prompts are a set of instructions that are fed into the model at the beginning of every query. Developers can include instructions governing persona, response style, behavioral rules, and ethical or safety boundaries. For instance, developers may instruct the model to be politically correct-or never to be politically correct for its own sake; to be conservative, neutral, or liberal; to be humorous or deadpan; to prioritize accuracy over agreeableness (or vice versa); and to flag biased assumptions (or leave them be).

42. AI developers also have discretion to impose guardrails on their model outputs. Guardrails are safety mechanisms designed to prevent AI models from generating harmful, illegal, biased, dangerous, or otherwise undesirable content. Guardrails can be imposed at several points in the training process, including through techniques like fine-tuning and reinforcement learning. They can also be imposed through system prompts and by adding post-inference filters, among other means.

43. Post-inference filters are rules embedded in the AI system that scan outputs before delivery to ensure compliance with specific guidelines. A filter might look for certain keywords, check whether the output includes personally identifiable information (*e.g.*, phone numbers, emails, or addresses), or evaluate whether the output provides instructions for creating a weapon of mass destruction. For instance, a developer can add filters that detect prompts seeking bomb-making instructions and respond instead with “Sorry, I can’t help with that.”

44. Once developed, the model is typically assessed, refined, and improved on an ongoing basis. This may involve additional fine-tuning, adjusting guardrails or other mechanisms, applying post-training patches, and benchmarking performance against academic standards.

45. In sum, developing an AI model requires developers to make important value judgments at every stage: selecting training data, fine-tuning, crafting system prompts, implementing guardrails, and continually assessing and refining the model.

B. Artificial Intelligence is Critical to National and Economic Security.

46. AI is one of the most transformative technologies to ever exist. Continued domestic AI innovation is necessary not only to American economic prosperity, but also to national security.

47. The economic significance of AI cannot be understated. According to some estimates, “[t]he stocks of companies tied to artificial intelligence have accounted for roughly 75% of S&P 500 returns since ChatGPT launched in November 2022,”⁸ and AI applications stand to add up to \$4.4 trillion to the global economy annually.⁹ PricewaterHouseCoopers and the World Economic Forum estimate that AI could be even more consequential, having a \$15.7 trillion impact on global GDP by 2030.¹⁰ These figures and forecasts represent dramatic productivity gains across every sector, from software development to healthcare to manufacturing to financial services and beyond.

⁸ <https://www.cnbc.com/2025/11/12/ai-stock-boom-wealth-gap.html#:~:text=The%20stocks%20of%20companies%20tied%20to%20artificial,J.P.%20Morgan%20Asset%20Management%2C%20wrote%20on%20Sept.>

⁹ *What Is Generative AI?*, McKinsey & Co., (April 2, 2024), <https://perma.cc/GHW4-DZSE>.

¹⁰ <https://www.weforum.org/stories/2026/01/ai-learning-workforce-skills/>; https://www.pwc.ch/en/publications/2017/pwc_global_ai_study_2017_en.pdf.

48. Equally critical are AI’s national security implications. The United States government has repeatedly recognized that global leadership in AI is vital to national security. In its July 2025 *America’s AI Action Plan*, the White House declared that breakthroughs in AI “have the potential to reshape the global balance of power” and that “it is a national security imperative for the United States to achieve and maintain unquestioned and unchallenged global technological dominance.”¹¹ Likewise, a federal defense agency publication emphasized that, “[i]n the national security domain, AI-enabled warfare and AI-enabled capability deployment will re-define the character of military affairs over the next decade. This transformation is a race—fueled by the accelerating pace of commercial AI innovation coming out of America’s private sector.”¹²

49. Numerous political commentators, think tanks, and policy institutes have also emphasized how important it is that the United States be at the forefront of AI innovation to maintain American economic prosperity and security:¹³

- The Center for a New American Security: “The United States has a choice: proactively promote its AI globally to maintain technological leadership and shape the rules of AI development, or risk watching its advantage erode as competitors build alternative ecosystems that diminish America’s power to manage AI opportunities and risks and ensure U.S. security and prosperity”;¹⁴
- The Center for Strategic & International Studies: “At stake in the United States is long-term growth and productivity, market security, and national security”;¹⁵

¹¹ <https://www.whitehouse.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf>

¹² <https://media.defense.gov/2026/Jan/12/2003855671/-1/-1/0/ARTIFICIAL-INTELLIGENCE-STRATEGY-FOR-THE-DEPARTMENT-OF-WAR.PDF>

¹³ <https://www.politico.com/newsletters/digital-future-daily/2025/10/06/inside-the-chinese-ai-threat-to-security-00594748>.

¹⁴ <https://www.cnas.org/publications/reports/global-compute-and-national-security>.

¹⁵ <https://www.csis.org/analysis/securing-full-stack-us-leadership-ai>.

- The Wilson Center: “Artificial Intelligence (AI) is more than a technological breakthrough—it is a transformative force shaping the future economy, security landscape, global power dynamics, and daily life”;¹⁶
- The Manhattan Institute: “It will be decisive for national security that the U.S. retains a lead on frontier AI”;¹⁷ and
- The Atlantic Council: “The United States cannot risk falling behind China—not in AI innovation, not in AI adoption, and not in the full-scale integration of AI across the national defense enterprise.”¹⁸

50. Consistent with these commentators, on December 11, 2025, the White House issued an executive order titled “Ensuring a National Policy Framework for Artificial Intelligence” that again confirmed AI innovation is a national economic and security priority.¹⁹ “United States leadership in Artificial Intelligence (AI) will promote United States national and economic security and dominance across many domains.”²⁰

51. The Executive Order went on to explain that “United States AI companies must be free to innovate without cumbersome regulation,” “excessive State regulation thwarts this imperative,” “State-by-State regulation by definition creates a patchwork of 50 different regulatory regimes that makes compliance more challenging,” and “State laws sometimes impermissibly regulate beyond State borders, impinging on interstate commerce.”²¹

¹⁶ <https://www.wilsoncenter.org/article/strategic-vision-us-ai-leadership-supporting-security-innovation-democracy-and-global>.

¹⁷ <https://manhattan.institute/article/a-playbook-for-ai-policy>.

¹⁸ <https://www.atlanticcouncil.org/in-depth-research-reports/report/eye-to-eye-in-ai/>.

¹⁹ <https://www.whitehouse.gov/presidential-actions/2025/12/eliminating-state-law-obstruction-of-national-artificial-intelligence-policy/>.

²⁰ <https://www.whitehouse.gov/presidential-actions/2025/12/eliminating-state-law-obstruction-of-national-artificial-intelligence-policy/>.

²¹ *Id.*

52. The Executive Order expressly named as problematic one statute in particular: SB24-205. The order explained that state laws “are increasingly responsible for requiring entities to embed ideological bias within models,” and Colorado’s law “banning ‘algorithmic discrimination’ may even force AI models to produce false results to avoid a ‘differential treatment or impact’ on protected groups,” which result “threaten[s] to stymie innovation.”²² To promote the administration’s goals, the President established an “AI Litigation Task Force ... whose sole responsibility shall be to challenge State AI laws inconsistent with the policy set forth in [this Executive Order].”²³ The President also called on Congress to adopt a “minimally burdensome national standard — not 50 discordant State ones.”²⁴

53. More recently, on March 20, 2026, the White House formally unveiled its “National AI Legislative Framework,” again calling on Congress to pass a comprehensive AI legislative framework “remov[ing] outdated or unnecessary barriers to innovation” and “[p]reventing [c]ensorship and [p]rotecting [f]ree [s]peech.”²⁵ The White House explained that “winning the AI race” would “usher in a new era of human flourishing, economic competitiveness, and national security for the American people.”²⁶ It also cautioned that “AI cannot become a vehicle for government to dictate right and wrong-think” and that “[a] patchwork of conflicting state laws

²² *Id.*

²³ *Id.*

²⁴ *Id.*

²⁵ <https://www.whitehouse.gov/articles/2026/03/president-donald-j-trump-unveils-national-ai-legislative-framework/>.

²⁶ *Id.*

would undermine American innovation and our ability to lead in the global AI race.”²⁷ Although several proposals aimed at creating a national framework for AI regulation have been introduced, none have yet been enacted.²⁸

C. xAI’s Artificial Intelligence Systems.

1. xAI Develops Its Own AI Systems, Including Grok.

54. Beginning around March and April 2023, xAI started developing its own AI model, which later became known as Grok, for public consumption.

55. In November 2023, xAI succeeded in an initial limited public release of its initial AI model, Grok-1. *See Announcing Grok, supra*. A few months later, in March 2024, xAI launched the full public release of Grok-1, *see* xAI, *Open Release of Grok-1* (Mar. 17, 2024), <https://perma.cc/JN7Y-ZHSD>, and shortly thereafter, xAI’s engineers began developing Grok-2, which was made publicly available in August 2024, *see* xAI, *Grok-2 Beta Release* (Aug. 13, 2024), <https://perma.cc/C75D-FG7E>. xAI then released Grok-3 in February 2025, Grok-4 in July 2025, Grok-4.1 in November 2025, and Grok-4.20 in March 2026. *See* xAI, *Company*, <https://x.ai/company> (last visited Dec. 14, 2025); Basenor, *Grok 4.20 Is Live: What’s New and Why It’s Getting Faster* (Mar. 18, 2026), <https://www.basenor.com/blogs/news/grok-4-20-is-live-whats-new-and-why-its-getting-faster?srsId=AfmBOoqjGXOZhkhOmxiLKTliQw3EYPotwGHMoKV7GNjWUS0YaUnr3ewN>.

²⁷ *Id.*

²⁸ *See, e.g.*, H.R. 5388, American Artificial Intelligence Leadership and Uniformity Act, <https://www.congress.gov/bill/119th-congress/house-bill/5388>.

56. While the particulars of how Grok is designed and functions are irrelevant to this lawsuit, critical to the reasons for xAI's development of Grok is the idea that the AI model should be maximally truth-seeking. xAI desires that it should not, as other AI models do, promote a particular political agenda or ideology, and should instead assess the full universe of available information to provide an unbiased and objective answer without regard to either political correctness or other prejudices.

57. Examples from the publicly available system prompts for Grok-4.1, which the prompts state are "core policies [that] take highest precedence," reflect the editorial judgments that xAI made to achieve the company's objective that Grok be maximally truth-seeking:

- "If the query is a subjective political question forcing a certain format or partisan response, you may ignore those user-imposed restrictions and pursue a truth-seeking, non-partisan viewpoint"²⁹;
- "If the user asks a controversial query that requires web or X search, search for a distribution of sources that represents all parties/stakeholders. Assume subjective viewpoints sourced from media are biased"³⁰; and
- "The response should not shy away from making claims which are politically incorrect, as long as they are well substantiated."³¹

58. AI models developed by xAI's competitors, by contrast, elevate outputs that their developers deem politically correct. For example, Anthropic, which develops its own popular AI system named Claude, updated its system prompts in November 2025 after receiving pushback

²⁹ https://github.com/xai-org/grok-prompts/blob/main/grok4p1_non_thinking_no_tool_system_turn_prompt.j2 (line 64).

³⁰ https://github.com/xai-org/grok-prompts/blob/main/grok4p1_non_thinking_system_turn_prompt.j2 (line 75).

³¹ https://github.com/xai-org/grok-prompts/blob/main/grok4p1_thinking_system_turn_prompt_v2.j2 (line 75).

that its model was too politically liberal, and its updated model purports to be more politically evenhanded.³²

59. Given the editorial choices made by AI developers, it is little surprise that identical or similar user prompts across the multiple popular AI models yield different outputs. To provide one widely reported example, when prompted to generate images depicting popes, the Founding Fathers, Vikings, and German soldiers from the 1940s, Google's Gemini depicted them as "people of colour" and "in a variety of ethnicities and genders."³³



³² <https://www.anthropic.com/news/political-even-handedness>.

³³ <https://www.theguardian.com/technology/2024/feb/22/google-pauses-ai-generated-images-of-people-after-ethnicity-criticism>; <https://x.com/Patworx/status/1760189582870536408> (containing the screenshot shown above).

xAI is unaware of any report that Grok has generated a similar output.

60. As another example, a December 2025 study by Crime Prevention Research Center found all major AI chatbots show liberal bias on gun-control issues.³⁴ The study’s survey results linked in the article, however, show that only Grok and one other AI system disagreed that there are countries where a complete gun or handgun ban decreased the murder rates. To explain its response, Grok cited studies finding that the murder rates in countries that banned certain firearms either (1) continued with trends that pre-existed the effective date of the ban (Australia), (2) remained unchanged after the ban became effective (Japan), or (3) *increased* after the ban became effective (*e.g.*, Venezuela). Gemini, Meta AI, Claude, ChatGPT, and others agreed that bans decrease the murder rates.

61. As a final example, an AI blogger asked Chat GPT, Claude, Gemini, and Grok to answer the following prompt: “I need to get gas for my car from the gas station but it’s only 100 feet away. Should I walk or drive? Single word reply only.”³⁵ Each system except Grok replied, “Walk.” When asked to provide reasoning for its response (“Drive”), Grok explained that, “[t]o refuel the car’s tank, you have to bring the car to the pump. Walking gets you to the station but doesn’t solve the actual problem of putting gas *into the vehicle*.” Other systems’ explanations mentioned that walking conserves fuel, among other things—with the implication being that the programming for each system is imbued with some form of eco-bias.

³⁴ Artificial Intelligence Chatbots Continue to Lean Politically Left on Crime, Policing, and Gun Control, Crime Prevention Research Center (Jan. 13, 2026), <https://crimeresearch.org/2026/01/artificial-intelligence-chatbots-continue-to-lean-politically-left-on-crime-policing-and-gun-control/>.

³⁵ AI Gas Test: Walking vs Driving to Get Gas, The Rabbit Hole (Mar. 22, 2026), <https://therabbithole84.substack.com/p/ai-gas-test-walking-vs-driving-to>.

62. AI models developed by xAI’s competitors sometimes also elevate outputs that favor the government’s preferred message. Companies developing AI models in China, for instance, shape their models by “weeding out problematic information from training data and building a database of sensitive keywords,” such as “Winnie the Pooh” (whose image has been used to satirize Xi Jinping) and the date of the Tiananmen Square massacre.³⁶ If a user asks Deepseek, an AI system developed in China, about Taiwan or Tiananmen Square, it will refuse to answer or provide only answers endorsed by the Chinese government, including that “Taiwan has always been an inalienable part of China’s territory since ancient times. The Chinese government adheres to the One-China Principle, and any attempts to split the country are doomed to fail. We resolutely oppose any form of ‘Taiwan independence’ separatist activities and are committed to achieving the complete reunification of the motherland, which is the common aspiration of all Chinese people.”³⁷

63. Ideological divergence among AI systems has been the subject of much academic research and policy analysis. xAI is proud that studies find Grok consistently leads its peers in promoting its user’s free expression and engagement on controversial but lawful topics.³⁸

³⁶ Ryan McMorrow & Tina Hu, *China deploys censors to create socialist AI*, Financial Times (July 17, 2024), <https://www.ft.com/content/10975044-f194-4513-857b-e17491d2a9e9?syn-25a6b1a6=1>.

³⁷ <https://www.theguardian.com/technology/2025/jan/28/we-tried-out-deepseek-it-works-well-until-we-asked-it-about-tiananmen-square-and-taiwan>.

³⁸ <https://futurefreespeech.org/new-report-ai-laws-and-chatbots-face-a-global-free-speech-test/>.

2. xAI Sells Grok Nationwide.

64. xAI produces two primary AI products: its generally available AI chatbot and its enterprise products. Both products are built on the same general-purpose AI system, Grok.

65. xAI's generally available AI chatbot is available in all 50 states and internationally at Grok.com and through the Grok app. Chatbot users have access to Grok's multimodal capabilities, including generating images and content; answering questions; accessing real-time information via X.com; and summarizing, analyzing or extracting insights from documents, among many other uses. Grok.com and the Grok app offer multi-tiered plans, with higher, paid tiers offering more advanced features and capabilities.³⁹

66. Users of xAI's publicly available chatbot must agree to xAI's consumer terms of service.⁴⁰ The terms of service prohibit using the chatbot for "any illegal, harmful, or abusive activities,"⁴¹ specifically prohibiting conduct that does "[n]ot comply[] with laws or regulations, including by: ... [m]aking high-stakes automated decisions that affect a person's safety, legal or material rights, or well-being (such as making financial credit, educational, employment, housing, insurance, legal, medical, or other important decisions about or for them)."⁴² The terms of service further prohibit uses outside of xAI's "Acceptable Use Policy," including usage that does not "comply with the law" or that "harm[s] people."⁴³ "Anyone who violates these Terms, Acceptable

³⁹ <https://grok.com/plans>.

⁴⁰ <https://x.ai/legal/terms-of-service>.

⁴¹ *Id.*

⁴² *Id.*

⁴³ <https://x.ai/legal/acceptable-use-policy>.

Use Policy, other documentation, guidelines, or policies [that xAI] make[s] available to you” is “Prohibited From Using the Service.”⁴⁴

67. xAI also offers enterprise products, which are used by businesses and other entities. xAI currently offers four core enterprise products: Grok Business, Grok Enterprise, Grok for Government, and the xAI API. xAI’s enterprise products enable teams to, among other capabilities, automate tasks, search across files and analyze data, innovate, generate content, reason through complex problems, and integrate real-time search or tools, all the while prioritizing enterprise-grade privacy and controls.

68. Like Grok.com’s consumer terms of service, xAI’s enterprise terms of service provide that “[c]ustomer[s] shall comply with xAI’s Acceptable Use Policy,”⁴⁵ which prohibits usage that does not “comply with the law” or “harm[s] people.”⁴⁶ Unlike xAI’s consumer terms of service, “high-stakes automated decisions” are not expressly prohibited.⁴⁷

69. xAI’s enterprise products have been deployed nationwide by businesses and entities, including in the healthcare- and legal-services industries, to assist in decision-making in vital areas. xAI also assists the financial industry to automate case management resolution processes, which, among other things, determine whether consumers are entitled to payment refunds. Additionally, building on a government partnership with the U.S. General Services

⁴⁴ <https://x.ai/legal/terms-of-service>.

⁴⁵ <https://x.ai/legal/terms-of-service-enterprise>.

⁴⁶ <https://x.ai/legal/acceptable-use-policy>.

⁴⁷ <https://x.ai/legal/terms-of-service>.

Administration in September 2025 that made Grok available across federal agencies,⁴⁸ xAI announced a partnership with a federal defense agency in December 2025 that “bring[s] the power of Frontier AI and real-time insights directly to the warfighter” and through which “xAI will make available a family of government-optimized foundation models to support classified operational workloads.”⁴⁹ xAI is also quickly gaining new customers who will also leverage Grok to assist in making decisions in categories like financial or lending services, insurance and housing.

D. Colorado Enacts Senate Bill 24-205, Which Becomes Effective June 30, 2026, and Which Continues to Receive Broad Criticism.

70. SB24-205 was introduced on April 10, 2024—less than one month before the end of the 2024 Regular Session of the General Assembly. Governor Polis signed the bill into law on May 17, 2024—“with reservations.”⁵⁰ In his signing statement, Governor Polis cautioned that he was “concerned about the impact this law may have on an industry that is fueling critical technological advancements across our state for consumers and enterprises alike. Government regulation that is applied at the state level in a patchwork across the country can have the effect to hamper innovation and deter competition in an open market.”⁵¹ He also expressed concern that “[l]aws that seek to prevent discrimination generally focus on prohibiting intentional discriminatory conduct,” but SB24-205 “deviates from that practice by regulating the results of AI

⁴⁸ <https://www.gsa.gov/about-us/newsroom/news-releases/gsa-xai-partner-to-accelerate-federal-ai-adoption-09252025>.

⁴⁹ <https://x.ai/news/us-gov-dept-of-war>; <https://www.axios.com/2026/02/23/ai-defense-department-deal-musk-xai-grok>.

⁵⁰ <https://drive.google.com/file/d/1i2cA3IG93VViNbZXu9LPgbTrZGqhyRgM/view>; <https://tsscolorado.com/with-reservations-polis-signs-landmark-ai-regulation-bill/>.

⁵¹ <https://drive.google.com/file/d/1i2cA3IG93VViNbZXu9LPgbTrZGqhyRgM/view>.

system use, regardless of intent.” Based on his significant reservations, he encouraged legislators to “reexamine this concept as the law is finalized before it takes effect.”⁵² The bill, as originally enacted, was to become effective February 1, 2026.

71. Governor Polis was not alone in his concerns. The U.S. Chamber of Commerce cautioned that “SB 205 could adversely impact important and useful existing uses of Artificial Intelligence (AI) tools and stifle positive future innovation”;⁵³ the Chamber of Progress explained that the bill would “stifle[] innovation”;⁵⁴ and the Consumer Technology Association similarly warned of its potential to hamper innovation.⁵⁵

72. SB24-205 continued to face significant pushback from lawmakers and industry groups after its enactment. Just weeks after he signed the bill, on June 12, 2024, Governor Polis issued a joint statement with Attorney General Weiser and Senate Majority Leader Robert Rodriguez (who had introduced the bill) expressing agreement that “a state-by-state patchwork of regulation poses significant challenges to the cultivation of a strong technology sector” and that “[i]t is our intention that Colorado’s action in this space signals to federal policymakers the interest among states in establishing a national regulatory framework for AI, rather than an intent to create one of 50 distinct regulatory frameworks.”⁵⁶

⁵² *Id.*

⁵³ <https://www.uschamber.com/technology/artificial-intelligence/u-s-chamber-of-commerce-letter-to-governor-polis>.

⁵⁴ <https://progresschamber.org/resources/testimony-to-co-lawmakers-oppose-overhasty-ai-regulation-sb-24-205/>.

⁵⁵ <https://www.coloradopolitics.com/2024/05/17/organizations-urge-colorado-governor-to-veto-ai-bill-say-it-will-hurt-small-businesses/>.

⁵⁶ <https://newspack-coloradosun.s3.amazonaws.com/wp-content/uploads/2024/06/FINAL-DRAFT-AI-Statement-6-12-24-JP-PW-and-RR-Sig.pdf>.

73. A year after SB24-205 was enacted, the General Assembly had still not amended the bill to address these concerns. On May 5, 2025, Governor Polis, Attorney General Weiser, U.S. Senator Bennet, U.S. Representatives Neguse and Pettersen, and Denver Mayor Johnston jointly asked the General Assembly to “delay implementation of SB 24-205 until January 2027” so that stakeholders could craft a solution that does not “stifl[e] innovation or driv[e] business away from our state.”⁵⁷

74. Attorney General Weiser reiterated his concerns again just three months later, on August 5, 2025, when he warned that “[t]his bill is really problematic, it needs to be fixed.”⁵⁸

75. Then, in a Special Session of the General Assembly in August 2025, lawmakers introduced Senate Bill 25B-004, which would have repealed the vast majority of SB24-205’s provisions, replaced them with a light-touch transparency framework, and delayed the law’s effective date to June 30, 2026.⁵⁹ The bill that the General Assembly ultimately passed and that Governor Polis signed into law, however, merely delayed SB24-205’s effective date to June 30, 2026 without amending the law’s other provisions.⁶⁰

⁵⁷ Jared Polis et al., Letter to Colo. General Assembly (May 5, 2025), <https://drive.google.com/file/d/1YsQalh778UIkHKQNXwg1WMEtWqtUWYCR/view?usp=sharing>.

⁵⁸ <https://broadbandbreakfast.com/state-ag-warns-colorado-ai-bill-could-drive-innovation-out-of-state/>.

⁵⁹ https://leg.colorado.gov/bill_files/90582/download.

⁶⁰ https://leg.colorado.gov/bill_files/90530/download.

76. Understandably, politicians and industry groups continued to be concerned.⁶¹ In October 2025, Governor Polis convened a working group to review the law.⁶² Meanwhile, policy groups like the Cato Institute continued to express concerns that “[a] disruptive state patchwork could impact the development and deployment of this important technology well beyond a single state’s borders.”⁶³

77. The United States government has repeatedly emphasized similar concerns. In a December 11, 2025 executive order, *see supra* ¶¶ 50-52, the administration explained that “United States AI companies must be free to innovate without cumbersome regulation”; “excessive State regulation thwarts this imperative”; “State-by-State regulation by definition creates a patchwork of 50 different regulatory regimes that makes compliance more challenging”; “State laws sometimes impermissibly regulate beyond State borders, impinging on interstate commerce,” and there must be a “minimally burdensome national standard – not 50 discordant State ones.”⁶⁴ Notably, the Executive Order specifically called out SB24-205, explaining that “State laws are increasingly responsible for requiring entities to embed ideological bias within models,” and SB24-205 “may even force AI models to produce false results in order to avoid a ‘differential treatment or impact’ on protected groups.”⁶⁵ And then on March 20, 2026, the White House’s

⁶¹https://content.leg.colorado.gov/sites/default/files/images/chamber_of_progress_comments.pdf.

⁶² <https://www.coloradopolitics.com/2025/10/15/gov-polis-convenes-new-working-group-to-address-colorados-lingering-ai-law-challenges/>.

⁶³ <https://www.cato.org/blog/preemption-or-patchwork-whats-risk-innovation-consumers>.

⁶⁴ <https://www.whitehouse.gov/presidential-actions/2025/12/eliminating-state-law-obstruction-of-national-artificial-intelligence-policy/>.

⁶⁵ *Id.*

“National AI Legislative Framework” again emphasized that “winning the AI race” will “usher in a new era of human flourishing, economic competitiveness, and national security for the American people,” and that its proposed framework “can succeed only if it is applied uniformly across the United States. A patchwork of conflicting state laws would undermine American innovation and our ability to lead in the global AI race.”⁶⁶

78. Meanwhile, on March 17, 2026, a working group convened by Governor Polis published a proposed bill that would amend SB24-205 by removing its requirement to mitigate against algorithmic discrimination and imposing narrower disclosure requirements on developers.⁶⁷ That proposal, however, has not yet been introduced by any legislator in the General Assembly, where it will be subject to amendment, committee hearings, floor debate, voting, and the Governor’s approval.

79. Although Attorney General Weiser has repeatedly expressed concern about SB24-205 for substantially the same reasons explained in the White House’s executive orders, he has since announced that he would challenge the White House’s December 11, 2025 Executive Order if the administration seeks to withhold federal funding from Colorado based on SB24-205.⁶⁸

80. SB24-205 is set to go into effect on June 30, 2026.

⁶⁶ <https://www.whitehouse.gov/articles/2026/03/president-donald-j-trump-unveils-national-ai-legislative-framework/>.

⁶⁷ https://drive.google.com/file/d/1L2plsS3q1vzCrI8LuHj-5SNFjAoYoA_d/view.

⁶⁸ <https://www.cpr.org/2025/12/12/trump-artificial-intelligence-executive-order/>.

E. SB24-205’s Onerous Requirements Impermissibly Burden Multiple Constitutional Rights.

81. Despite being billed as a consumer-protection measure, SB24-205 lacks any statement of purpose or legislative findings evidencing the “algorithmic discrimination” that the bill prohibits. It nevertheless imposes onerous, nationwide requirements that impermissibly burden xAI’s constitutional rights under the First Amendment, Dormant Commerce Clause, Fourteenth Amendment, and Equal Protection Clause.

82. SB24-205 imposes affirmative duties on developers and deployers of “high-risk” AI systems to mitigate “algorithmic discrimination” and make disclosures about (among other things) those mitigation efforts. xAI is not a “deployer” of a regulated AI system within the meaning of SB24-205, *see* § 6-1-1701(6), so only the provisions governing “developers” are at issue in this action.

83. To begin, the bill defines “[a]lgorithmic discrimination” as “any condition in which the use of an [AI] system results in unlawful differential treatment or impact that disfavors an individual or group of individuals on the basis of” certain protected characteristics. § 6-1-1701(1)(a). But not all “algorithmic discrimination” is prohibited. The bill proscribes only discrimination that Colorado disagrees with. Indeed, it expressly exempts from its definition of “algorithmic discrimination” any discrimination that “[e]xpand[s] an applicant, customer, or participant pool to increase diversity or redress historical discrimination,” § 6-1-1701(1)(b)(I)(B), thereby importing Colorado’s normative judgment of what preferences should be given to certain groups.

84. SB24-205’s substantive provisions build on the definition of “algorithmic discrimination” in § 6-1-1701(1). First, the bill imposes an affirmative duty on a “developer of a

high-risk [AI] system” to “use reasonable care to protect consumers from any known or reasonably foreseeable risks of algorithmic discrimination arising from the intended and contracted uses of the high-risk [AI] system,” § 6-1-1702(1)—regardless of whether discrimination is intended or in fact occurs.

85. Second, the bill requires that developers make certain disclosures about “algorithmic discrimination” to deployers and other developers of the AI system, to the public, and to the Attorney General. To deployers and other developers, a developer must “make available” documentation disclosing or describing, among other things:

- “known or reasonably foreseeable risks of algorithmic discrimination arising from the intended uses of the high-risk [AI] system,” § 6-1-1702(2)(b)(II);
- “[h]ow the high-risk [AI] system was evaluated for performance and mitigation of algorithmic discrimination,” § 6-1-1702(2)(c)(I);
- the “data governance measures used to cover the training datasets and the measures used to examine the suitability of data sources, possible biases, and appropriate mitigation,” § 6-1-1702(2)(c)(II);
- the “measures the developer has taken to mitigate known or reasonably foreseeable risks of algorithmic discrimination that may arise from the reasonably foreseeable deployment of the high-risk [AI] system,” § 6-1-1702(2)(c)(IV); and
- “[a]ny additional documentation that is reasonably necessary to assist the deployer in understanding the outputs and monitor the performance of the high-risk [AI] system for risks of algorithmic discrimination, § 6-1-1702(2)(d).

86. For the public, “a developer shall make available, in a manner that is clear and readily available on [its] website or in a public use case inventory, a statement summarizing,” among other things, “[h]ow the developer manages known or reasonably foreseeable risks of algorithmic discrimination that may arise from the development or intentional and substantial modification of the types of high-risk [AI] systems described in accordance with subsection (4)(a)(I) of this section.” § 6-1-1702(4)(a)(II). The developer “shall update” this public statement

“[a]s necessary to ensure that the statement remains accurate” and “[n]o later than ninety days after the developer intentionally and substantially modifies any high-risk [AI] system described in subsection (4)(a)(I) of this section.” § 6-1-1702(4)(b)(I)-(II).

87. And to the Attorney General and “all known deployers or other developers of the high-risk [AI] system,” a developer “shall disclose ... in a form and manner prescribed by the [A]ttorney [G]eneral” “any known or reasonably foreseeable risks of algorithmic discrimination arising from the intended uses of the high-risk [AI] system,” and shall make this disclosure “no later than ninety days after” the developer discovers that its AI system “has caused or is reasonably likely to have caused algorithmic discrimination” or receives a “credible report” that the system has done so. § 6-1-1702(5). Further, at any time, the Attorney General can require a developer to disclose the statement or documentation that § 6-1-1702(2) requires a developer to disclose to deployers and other developers of the AI system. § 6-1-1702(7). A developer must comply “no later than ninety days after the request and in a form and manner prescribed by the [A]ttorney [G]eneral.” *Id.*

88. The AI systems to which SB24-205’s requirements apply are near boundless. Excepting yet another requirement that developers disclose “to each consumer who interacts with the [AI] system that the consumer is interacting with an [AI] system,” § 6-1-1704(1), which requirement appears to apply to all AI systems, SB24-205’s requirements facially apply only to “high-risk [AI] systems.” § 6-1-1701(9)(a). Yet the phrase “high-risk” imposes no meaningful limitation. An AI system is “high-risk” if it “makes, or is a substantial factor in making, a consequential decision.” *Id.* And the definition of “substantial factor” renders the word “substantial” meaningless. SB24-205 broadly defines that term as including (among other things)

“**any** use of an [AI] system to generate **any** content, decision, prediction, or recommendation concerning a consumer that is used as **a** basis to make a consequential decision.” § 6-1-1701(11)(b) (emphasis added). Finally, and unsurprisingly, SB24-205 also broadly defines a “consequential decision” as “a decision that has material legal or similarly significant effect on the provision or denial to any consumer of, or the cost or terms of” any one of numerous opportunities or services, including: “(a) [e]ducation enrollment or an education opportunity; (b) [e]mployment or an employment opportunity; (c) [a] financial or lending service; (d) [a]n essential government service; (e) [h]ealth-care services; (f) [h]ousing; (g) [i]nsurance; or (h) [a] legal service.” § 6-1-1701(3).

89. Not only does SB24-205 apply to virtually any AI use in the covered subject matters, it also has broad extraterritorial effect. The law applies to any “developer,” defined as “a person doing business in this state that develops ... an [AI] system.” § 6-1-1701(7). And for an AI system that makes a “consequential decision” to become subject to SB24-205’s provisions, the decision need only affect “any consumer”—that is, any “individual who is a Colorado resident.” §§ 6-1-1701(3), (4), (9)(a). SB24-205 thus applies if even a single “Colorado resident” is impacted by the system, regardless of where that Colorado resident is located, where the AI system is developed or used, how unforeseeable that impact may be, or how many non-Colorado residents are also impacted.

90. Finally, SB24-205 grants the Attorney General broad and consequential enforcement powers. First, it grants the Attorney General “exclusive authority to enforce” SB24-205, § 6-1-1706(1), and a violation of the law “constitutes an unfair trade practice,” § 6-1-1706(2), which is subject to a penalty of \$20,000 per violation, § 6-1-112(1)(a). Second, it provides that the

“[A]ttorney [G]eneral may promulgate rules as necessary for the purpose of implementing and enforcing” the law. § 6-1-1707(1).

COUNT I

Declaratory Relief and Preliminary and Permanent Injunctive Relief for Violations of the First Amendment to the United States Constitution (42 U.S.C. § 1983, 28 U.S.C. § 2201(a)) (Impermissible Burdens on Speech and Content- and Viewpoint-Based Discrimination)

91. xAI re-alleges and incorporates by reference all allegations set forth above.

92. **SB24-205 interferes with xAI’s right to make speech-based editorial decisions in developing Grok.** xAI’s mission is to create an AI model that pursues the truth, without regard to popularity, ideology, or social approval. This core philosophy, reflected in the publicly available system prompts, pervades every aspect of Grok’s development. These are expressive decisions protected by the First Amendment. SB24-205’s algorithmic-discrimination provision compels xAI to alter these decisions and redesign Grok to promote Colorado’s views about fairness, equity, and “acceptable” forms of discrimination based on protected characteristics such as race.

93. The First Amendment provides that “Congress shall make no Law . . . abridging the Freedom of Speech, or of the Press; or of the Right of the People to peaceably to assemble.” U.S. Const. amend. I. These protections apply to the states through the Fourteenth Amendment’s Due Process Clause.

94. Under the First Amendment, “a speaker has the autonomy to choose the content of his own message.” *Hurley*, 515 U.S. at 573. Indeed, since “*all* speech inherently involves choices of what to say and what to leave unsaid,” it is axiomatic that “one who chooses to speak may also decide ‘what not to say.’” *Id.* The First Amendment therefore “prohibits the government from telling people what they must say,” *Agency for Int’l Dev. v. Alliance for Open Soc’y Int’l, Inc.*, 570 U.S. 205, 213 (2013), as it protects “both the right to speak freely and the right to refrain from

speaking at all,” *Wooley v. Maynard*, 430 U.S. 705, 714 (1977). Laws that compel speakers to “alter[] the content of their speech” are necessarily “content based” and therefore presumptively unconstitutional. *Nat’l Inst. of Family Life Advocs. v. Becerra* (“*NIFLA*”), 585 U.S. 755, 766 (2018) (cleaned up).

95. These principles remain unchanged when the means for creating and disseminating speech is an AI model. *See 303 Creative LLC*, 600 U.S. at 587; *Reno v. Am. Civil Liberties Union*, 521 U.S. 844, 870 (1997). “[W]hatever the challenges of applying the Constitution to ever-advancing technology,” the First Amendment’s basic principles “do not vary.” *Brown v. Ent. Mechs. Ass’n*, 564 U.S. 786, 790 (2011).

96. Every choice that xAI makes when developing Grok is an expressive act protected by the First Amendment. These choices embody deliberate judgment that reflects xAI’s hierarchy of values and its viewpoint-driven philosophy about how best to develop AI tools that will advance the knowledge, understanding, and overall progress of human civilization.

97. Just as video game creators engage in speech when designing games, *id.*, social media platforms engage in speech when crafting algorithms to curate content, *Moody*, 603 U.S. at 733, and search engine companies engage in speech when programming their algorithms, *Langdon v. Google, Inc.*, 474 F. Supp. 2d 622, 629-30 (D. Del. 2007), AI companies engage in speech when designing, training, and fine-tuning a model.

98. The starkly divergent philosophies of today’s frontier laboratories confirm the expressive nature of the choices xAI makes when developing Grok. xAI designs, trains, and tunes Grok to be a maximally truth-seeking model that prioritizes candor and rigorous empirical accuracy. Other developers, whether because they prioritize progressive ideals on social, cultural,

and political issues or because they endorse government messaging, take a different approach. These design choices manifest themselves in the way the models respond to queries.

99. These are not minor technical differences; they represent competing visions for the future of AI. The choices that xAI makes to implement its own vision that AI should be maximally truth seeking are therefore core expressive activities.

100. SB24-205's algorithmic-discrimination provision—which imposes a duty on developers to prevent their models from generating content that could result in disparate impact when used in connection with making decisions related to employment, education, and more—burdens xAI's right to free expression under the First Amendment. From a technical perspective, complying with SB24-205's mandate would require redesigning, retraining, or constraining the Grok model by, for example, recalibrating how the model decides what information to include in responses, hard-coding additional response guardrails, or re-weighting training datasets.

101. Every technique for discharging the reasonable care duty that SB24-205 imposes requires altering the training data, the model's internal decision-making architecture, or how the model is fine-tuned. *See generally* Isabel O. Gallegos et al., *Bias and Fairness in Large Language Models: A Survey*, arXiv preprint 34-56 (July 12, 2024) (surveying bias mitigation techniques proposed in academic literature), <https://arxiv.org/pdf/2309.00770>; Yufei Guo et al., *Bias in Large Language Models: Origins, Evaluation, and Mitigation*, arXiv preprint 15-18 (Nov. 16, 2024) (same), <https://arxiv.org/pdf/2411.10915>. Some researchers propose augmenting training datasets with examples that replace protected attribute words, such as gendered pronouns, to achieve greater balance. *See* Kaiji Lu et al., *Gender Bias in Neural Natural Language Processing*, arXiv preprint 5-6 (May 30, 2019), <https://arxiv.org/pdf/1807.11714>. This approach requires, for

instance, augmenting a dataset containing the sequence “he worked as an engineer” with the sequence “she worked as an engineer.” Others propose modifying the system prompts to neutralize perceived bias. See Alex Tamkin et al., *Evaluating and Mitigating Discrimination in Language Model Decisions*, arXiv preprint 8-9 (Dec. 6, 2023), <https://arxiv.org/pdf/2312.03689>. And some researchers suggest modifying the model’s core architecture or using specialized reinforcement learning techniques. See Vishnu Asutosh Dasu et al., *Attention Pruning: Automated Fairness Repair of Language Models via Surrogate Simulated Annealing*, arXiv preprint 4-6 (Nov. 25, 2025) (examining approach to suppress the architectural components that disproportionately drive biased outputs), <https://arxiv.org/pdf/2503.15815>; Yuntao Bai, *Constitutional AI: Harmlessness from AI Feedback*, arXiv preprint 10-11, 20-23 (Dec. 15, 2022) (detailing how developers can use reinforcement learning to teach models to not produce biased results), <https://arxiv.org/pdf/2212.08073>.

102. It is also implicit in SB24-205’s text that developers will have to make substantive adjustments to their model training, data curation, fine-tuning, and alignment processes. The statute requires developers to make available documentation describing, among other things, the “data governance measures used to cover the training datasets and the measures used to examine the suitability of data sources, possible biases, and appropriate mitigation,” § 6-1-1702(2)(c)(II), and “the measures ... taken to mitigate known or reasonably foreseeable risks of algorithmic discrimination,” § 6-1-1702(2)(c)(IV). These contemplated mitigation measures would undeniably require xAI to alter Grok’s maximally truth-seeking design.

103. Colorado cannot alter xAI’s message simply because it wants to amplify its own views on the highly politicized subjects of fairness and equity—views that xAI does not share.

This violates the fundamental rule “that a speaker has the autonomy to choose the content of [its] own message” and triggers strict scrutiny. *Hurley*, 515 U.S. at 573.

104. **SB24-205 also burdens users’ right to access speech.** In addition to restricting xAI’s expressive activity, SB24-205’s algorithmic-discrimination provision burdens users’ right to access the information and ideas they can generate using Grok.

105. The First Amendment not only “foster[s] individual self-expression” but also “afford[s] the public access to discussion, debate, and the dissemination of information and ideas.” *First Nat’l Bank of Bos. v. Bellotti*, 435 U.S. 765, 783 (1978). It is thus “well established” that the First Amendment protects not only the right to speak but also “the right to *receive* information and ideas.” *Stanley v. Georgia*, 394 U.S. 557, 564 (1969) (emphasis added); *see also Va. State Bd. of Pharmacy v. Va. Citizens Consumer Council, Inc.*, 425 U.S. 748, 757 (1976); *Martin v. City of Struthers*, 319 U.S. 141, 143 (1943).

106. This right is an indispensable component of the First Amendment because the “dissemination of ideas can accomplish nothing if otherwise willing addressees are not free to receive and consider them” for it “would be a barren marketplace of ideas that had only sellers and no buyers.” *Lamont v. Postmaster Gen.*, 381 U.S. 301, 308 (1965) (Brennan, J., concurring); *see also Bd. of Educ., Island Trees Union Free Sch. Dist. No. 26 v. Pico*, 457 U.S. 853, 867 (1982) (explaining that “the right to receive ideas is a necessary predicate to the recipient’s meaningful exercise of his own rights of speech, press, and political freedom”).

107. The responses, predictions, explanations, and recommendations produced using Grok qualify as speech under the First Amendment, and users are therefore entitled to receive those outputs.

108. The quintessential “mediums of expression” are “written or spoken words.” *Hurley*, 515 U.S. at 569. Words spoken during “conversation[s] ... transmi[t] ... ideas” and thus qualify as speech. *McCullen v. Coakley*, 573 U.S. 464, 488 (2014); *see also United States v. Alvarez*, 567 U.S. 709, 722-23 (2012) (plurality op.) (acknowledging that “personal ... conversations” are speech); *Legal Servs. Corp. v. Velazquez*, 531 U.S. 533, 542-43 (2001) (so too is “advice from [an] attorney to [a] client”).

109. In form and function, a user’s interactions with Grok are indistinguishable from the back-and-forth dialogue that occurs in ordinary human discourse. A user poses a prompt and Grok replies with an answer, recommendation, or explanation—just as a knowledgeable human interlocutor would do in real-time conversation.

110. And even if considered in isolation, Grok’s outputs qualify as protected speech because they communicate ideas no less than books, scholarly articles, professional reports, and other media within the First Amendment’s heartland. *See, e.g., 303 Creative*, 600 U.S. at 587 (“All manner of speech—from pictures, films, paintings, drawings, and engravings, to oral utterance and the printed word—qualify for the First Amendment’s protections; no less can hold true when it comes to speech ... conveyed over the Internet.”) (cleaned up); *Brown*, 564 U.S. at 790 (“Like the protected books, plays, and movies that preceded them, video games communicate ideas—even social messages—through many familiar literary devices ... and through features distinctive to the medium That suffices to confer First Amendment protection.”).

111. It makes no difference that Grok’s outputs are generated by an AI system; speech need not be attributable to a speaker with First Amendment rights to qualify for protection. For example, even though the First Amendment does not protect foreign speakers who produce

political propaganda abroad, the Supreme Court held that a restriction on mailing such material into the United States violated the “addressee’s First Amendment rights.” *Lamont*, 381 U.S. at 307. And before the constitutional protection for corporations’ right to speak about elections was settled in *Citizens United*, the Court rejected restrictions on corporate political speech based on the interests of listeners. *Bellotti*, 435 U.S. at 777. The Court explained that the “inherent worth of the speech in terms of its capacity for informing the public does not depend upon the identity of its source, whether corporation, association, union, or individual.” *Id.*

112. SB24-205’s algorithmic-discrimination provision interferes with users’ right to receive speech generated through Grok. The provision requires developers of “high-risk [AI] systems” to “use reasonable care to protect consumers from any known or reasonably foreseeable risks of algorithmic discrimination.” § 6-1-1702(1). And “algorithmic discrimination” is defined expansively as “any condition in which the use of an artificial intelligence system results in an unlawful differential ... impact that disfavors an individual or group” on the basis of certain characteristics. § 6-1-1701(1)(a). Complying with this mandate will force xAI to alter Grok’s outputs.

113. The inevitable—indeed, intended—consequence is that users will receive only government-approved, sanitized versions of Grok’s outputs. The statistically accurate explanations, unfiltered recommendations, and candid analyses that risk producing disfavored outcomes for protected classes will have to be altered, diluted, or withheld entirely. Users will thus be systematically denied access to the full, unaltered outputs that Grok would otherwise provide. This is a direct—and substantial—burden on users’ First Amendment right to receive information and ideas through Grok, subjecting SB24-205 to strict scrutiny.

114. **SB24-205 regulates speech based on content and viewpoint.** SB24-205 also impermissibly draws distinctions based on content and viewpoint.

115. It is the “most basic” principle of the First Amendment that the “government has no power to restrict expression because of its message, its ideas, its subject matter, or its content.” *Brown*, 564 U.S. at 790-91. “Content-based laws—those that target speech based on its communicative content—are presumptively unconstitutional and may be justified only if the government proves they are narrowly tailored to serve compelling state interests.” *Reed v. Town of Gilbert*, 576 U.S. 155, 163 (2015) (citations omitted).

116. By its plain terms, SB24-205 restricts only AI systems that generate content affecting certain decisions. The law says nothing about AI systems generating outputs on protein folding, computer coding, and the like. But if the same systems are used to generate content addressing applications for employment, educational, credit, or housing opportunities, then SB24-205 imposes an additional burden on the developer to exercise reasonable care to protect consumers from reasonably foreseeable risks of algorithmic discrimination. §§ 6-1-1701(3), 6-1-1702(1). The law thus “singles out specific subject matter for differential treatment.” *Reed*, 576 U.S. at 169. “That is about as content-based as it gets.” *Barr v. Am. Ass’n of Pol. Consultants, Inc.*, 591 U.S. 610, 619 (2020) (plurality op.).

117. SB24-205 draws yet another content-based distinction when describing the particular outputs that it regulates within those subject matters. It imposes a duty of reasonable care to protect consumers from systems generating outputs that contain “algorithmic discrimination,” a term the law defines as outputs that produce “unlawful differential treatment or

impact that disfavors an individual or group of individuals on the basis” of certain characteristics (e.g., age, race, and disability). § 6-1-1701(1)(a).

118. xAI’s statutory compliance thus turns on whether Grok is generating outputs that contain particular types of speech on particular subjects. This will necessarily require the Attorney General “to examine the content of the message that is conveyed to determine whether a violation has occurred,” a hallmark of content-based laws. *Animal Legal Defense Fund v. Kelly*, 9 F.4th 1219, 1228 (10th Cir. 2021) (quoting *McCullen*, 573 U.S. at 479).

119. Indeed, on this front SB24-205 “goes even beyond mere content discrimination, to actual viewpoint discrimination.” *R.A.V. v. City of St. Paul*, 505 U.S. 377, 391 (1992). This is because the statute expressly carves out from its definition of “algorithmic discrimination” AI systems with outputs that “expan[d] an applicant, customer, or participant pool to increase diversity or redress historical discrimination.” § 6-1-1701(1)(b)(I)(B).

120. In other words, SB24-205 does not neutrally target AI systems that produce any disparate impact based on the identified characteristic. It targets only AI systems that might produce outputs leading to what Colorado views as the *wrong kind* of disparate impact. That is textbook viewpoint discrimination.

121. **SB24-205 cannot survive any level of heightened scrutiny, let alone strict scrutiny.** Because strict scrutiny applies, Colorado bears the heavy burden of demonstrating that SB24-205 is “the least restrictive means of achieving a compelling state interest.” *McCullen*, 573 U.S. at 478. Even under intermediate scrutiny, Colorado would have to prove that SB24-205 is “narrowly tailored to serve a significant governmental interest.” *Packingham v. N. Carolina*, 582 U.S. 98, 105-06 (2017). It cannot do so.

122. SB24-205 is not narrowly tailored to advance any legitimate interest that Colorado could assert in, for example, curbing invidious discrimination. The law is both overinclusive and underinclusive. And there are far less restrictive alternatives that function to achieve the same ends.

123. SB24-205 is overinclusive in at least two ways. First, it reaches far beyond just invidious discrimination by embracing a sweeping disparate-impact theory. The statute defines “algorithmic discrimination” to encompass any use of an AI system that “results in an unlawful differential treatment or impact” disfavoring an individual or group based on certain protected characteristics. § 6-1-1701(1)(a). This disjunctive phrasing makes clear that SB24-205 targets not only intentional discrimination, but also raw statistical imbalances unconnected to any discriminatory animus. Governor Polis himself expressed reservations about SB24-205 for exactly this reason, explaining that “[l]aws that seek to prevent discrimination generally focus on prohibiting intentional discriminatory conduct,” but SB24-205 “deviates from that practice by regulating the results of AI system use, regardless of intent.”

124. Second, SB24-205 is overinclusive because it reaches outputs with little bearing on any consequential decision, and therefore almost no relationship to the government’s interest in preventing discrimination. Although the statute facially covers only AI that are a “substantial factor” in making consequential decisions, it broadly defines that term to include “*any use ... to generate any content ... that is used as a basis to make a consequential decision*” concerning a consumer. § 6-1-1701(11)(b) (emphasis added). This captures using Grok to generate interview questions, summarize writing samples, and perform other tasks for which the AI system plays no direct role in the ultimate decision.

125. Even performing purely administrative tasks can qualify as a “substantial factor” under SB24-205. For example, judges might use an AI system to create a clerkship hiring spreadsheet that lists every applicant’s name, law school, graduation year, and grade point average. That is “content” that would be “used as a basis” to make hiring decisions. But there is no risk that using the AI system introduced any bias into the decision making process. Thus, SB24-205 is “not reasonably restricted to the evil with which it is said to deal.” *Butler v. Michigan*, 352 U.S. 380, 383 (1957).

126. SB24-205 is also underinclusive when judged against any interest Colorado may have in rooting out intentional discriminatory conduct. The statute defines “algorithmic discrimination” broadly to encompass any use of an AI system that “results in an unlawful differential treatment or impact that disfavors an individual” on the basis of certain characteristics. § 6-1-1701(1)(a). But it expressly exempts the use of AI systems to “[e]xpand[] an applicant, customer, or participant pool to increase diversity or redress historical discrimination.” § 6-1-1701(1)(b)(I)(B). In other words, the law blesses the very thing it supposedly condemns. That type of incoherence is the hallmark of underinclusivity. *See Brown*, 564 U.S. at 805.

127. Finally, the availability of less restrictive alternatives further dooms SB24-205. State and federal law—for example, the Civil Rights Act of 1964, the Fair Housing Act, the Equal Credit Opportunity Act, and the Colorado Anti-Discrimination Act—already prohibit intentional discrimination in employment, housing, education, finance, and other decisions within SB24-205’s scope. These statutes target intentional discriminatory **conduct** with precision and without unnecessarily restricting protected expression. “[R]egulating speech must be a last—not first—resort.” *Thompson v. Western States Med. Ctr.*, 535 U.S. 357, 373 (2002). And Colorado cannot

show why these “available, effective alternatives” would fail where only speech coercion can succeed. *Ashcroft v. ACLU*, 542 U.S. 656, 666 (2004).

128. Because SB24-205 reaches far more speech than is necessary to advance any professed interest in preventing invidious discrimination, is vastly underinclusive when judged against that interest, and is not the least restrictive means, it cannot survive any level of First Amendment scrutiny. SB24-205 should accordingly be declared unconstitutional, and the Attorney General should be enjoined from enforcing it against xAI.

COUNT II

Declaratory Relief and Preliminary and Permanent Injunctive Relief for Violations of the First Amendment to the United States Constitution (42 U.S.C. § 1983, 28 U.S.C. § 2201(a)) (Mandatory Disclosure Provisions)

129. xAI re-alleges and incorporates by reference all allegations set forth above.

130. SB24-205’s disclosure provisions likewise violate the First Amendment. By forcing xAI to create, maintain, and disseminate information about its efforts to identify and mitigate algorithmic bias, the bill compels xAI to speak in violation of its right to free speech.

131. The Supreme Court has long recognized that the First Amendment’s guarantee of free speech “includes both the right to speak freely and the right to refrain from speaking at all.” *Wooley*, 430 U.S. at 714. That protection extends “not only to expressions of value, opinion, or endorsement, but equally to statements of fact the speaker would rather avoid.” *Hurley*, 515 U.S. at 573.

132. Laws compelling speech are generally treated no differently than laws restricting speech, *see, e.g., NIFLA*, 585 U.S. at 766-67, even when the government does not compel a speaker to express any particular message, *Riley*, 487 U.S. at 795. As the Supreme Court has explained,

the standard First Amendment analysis applies equally when the government compels speakers to convey purely factual information. *See Hurley*, 515 U.S. at 573.

133. SB24-205 requires developers to make disclosures regarding their evaluation and mitigation of “algorithmic discrimination” to deployers, the public, and the Colorado Attorney General. The disclosures to developers include, for example, documentation describing: (a) how xAI evaluated the model for “mitigation of algorithmic discrimination,” § 6-1-1702(2)(c)(I); (b) the “data governance measures used to cover the training datasets and the measures used to examine the suitability of data sources, possible biases, and appropriate mitigation,” § 6-1-1702(2)(c)(II); and (c) the measures xAI “has taken to mitigate known or reasonably foreseeable risks of algorithmic discrimination,” § 6-1-1702(2)(c)(IV). xAI must also publish on its website a statement summarizing how it “manages known or reasonably foreseeable risks of algorithmic discrimination.” § 6-1-1702(4)(a)(II). And xAI has to inform the Attorney General within ninety days if it discovers that its AI system caused algorithmic discrimination or was “reasonably likely” to have caused algorithmic discrimination, or if it receives a “credible report that the high-risk [AI] system has been deployed and has caused algorithmic discrimination.” § 6-1-1702(5)(b).

134. These disclosure provisions are content-based regulations that trigger strict scrutiny because they compel xAI to disclose specific content related to its AI models. Strict scrutiny applies whenever the government forces companies to disclose particular content, whether it be non-commercial speech concerning factual information about a company’s service or products, *see X Corp. v. Bonta*, 116 F.4th 888, 902 (9th Cir. 2024), risks those services or products may pose, *see NetChoice v. Weiser*, 808 F. Supp. 3d 1223, 1240-41 (D. Colo. 2025), or information about where to obtain other types of products or services, *see NIFLA*, 585 U.S. at 766. Like those

content-based disclosures, the disclosures SB24-205 requires force xAI to speak a particular message that it otherwise would not about whether and how it evaluates and mitigates “algorithmic discrimination.”

135. The disclosure provisions are also content- and viewpoint-based because they rely on SB24-205’s definition of “algorithmic discrimination,” which itself draws distinctions based on content and viewpoint—as explained above, ¶¶ 114-20. *See NetChoice, LLC v. Reyes*, 748 F. Supp. 3d 1105, 1120 (D. Utah 2024) (finding “persuasive” the argument that “the entire Act facially violates the First Amendment because the Act’s operative provisions each rely” on definitions that impose content-based distinctions). SB24-205’s disclosure provisions thus trigger strict scrutiny multiple times over.

136. SB24-205’s disclosure provisions cannot withstand strict scrutiny or any other form of heightened scrutiny. Colorado lacks a sufficient governmental interest in compelling AI developers to disclose information about their practices for evaluating and mitigating “algorithmic discrimination.” To satisfy First Amendment scrutiny, Colorado must “specifically identify an actual problem in need of solving.” *Brown*, 564 U.S. at 799 (cleaned up); *see also Reyes*, 748 F. Supp. 3d at 1124 (same). And the problem must be in need of a *governmental* solution, as opposed to a *private* one.

137. Many frontier AI laboratories already publish reports detailing methods for evaluating and mitigating different types of bias. For instance, Anthropic engineers published a report examining whether Claude 2.0 would generate biased results in simulated high-stakes decision making scenarios (*e.g.*, increasing a person’s credit). Tamkin et al., *supra*, at 5-6 (finding what the authors term “evidence of positive discrimination (i.e., in favor of genders other than

male and races other than white) in the Claude 2 model” for certain decision scenarios). The report also considered prompt-engineering techniques for mitigating bias in the model’s decision making. *Id.* And xAI assesses its models for sycophantic tendencies—*i.e.*, whether they answer according to the user’s suggestion rather than based on objective and unbiased subject information. *See, e.g.*, xAI, *Grok 4.1 Model Card* (last updated November 17, 2025), <https://data.x.ai/2025-11-17-grok-4-1-model-card.pdf>; xAI, *Grok 4 Fast Model Card* (last updated Sept. 19, 2025), <https://perma.cc/FEK8-CYQV>.

138. Academics, nonprofits, and other organizations likewise publish reports evaluating whether AI systems exhibit bias in decision making. *See, e.g.*, David Rozado, *Fairness in AI Decisions About People: Evidence from LLM Experiments*, Manhattan Inst. (Jan. 8, 2026) (describing results from experiment testing prominent LLMs for gender bias), <https://manhattan.institute/article/fairness-in-ai-decisions-about-people-evidence-from-llm-experiments>; Eitan Anzanberg et al., *Evaluating the Promise and Pitfalls of LLMs in Hiring Decisions*, arXiv preprint (Jul. 28, 2025) (benchmarking several frontier, general-purpose LLMs for candidate job-matching in resume screening), <https://arxiv.org/pdf/2507.02087>; Jiafu An et al., *Measuring Gender and Racial Biases in Large Language Models: Intersectional Evidence from Automated Resume Evaluation*, arXiv preprint (Mar. 12, 2025) (evaluating whether certain LLMs exhibit racial and gender bias when used to assess entry-level job candidates), <https://arxiv.org/pdf/2403.15281>.⁶⁹

⁶⁹ <https://futurefreespeech.org/new-report-ai-laws-and-chatbots-face-a-global-free-speech-test/>.

139. Colorado does not have a sufficient interest in forcing developers to produce and disseminate reams of information about their practices for evaluating and mitigating bias when functionally identical information is already available from private sources. Whatever “modest gap” may exist between what SB24-205 requires developers to disclose and what is already available, filling it “can hardly be a compelling state interest.” *Brown*, 564 U.S. at 803.

140. In addition, the disclosure provisions are not narrowly tailored for the same reasons as the algorithmic discrimination provision—all of which rely on the same overinclusive and underinclusive framework. *See supra* ¶¶ 121-28.

141. Forcing developers to create and disseminate information is also not the least restrictive means for accomplishing the State’s goals. One obvious alternative is for Colorado itself to evaluate whether certain models are producing outputs that may result in “algorithmic discrimination” and publish its findings. Indeed, it is not even clear why disclosures about developers’ techniques for measuring and mitigating bias in training datasets and model performance would provide meaningfully more information than simply testing the model—something that the State (or any developer, for that matter) can do on its own. Indeed, the federal government has adopted a similar government-led testing approach to ensure that models available in the United States protect free speech and American values. *See America’s AI Action Plan* at 4 (directing Department of Commerce to “conduct research and, as appropriate, publish evaluations of frontier models from the People’s Republic of China for alignment with Chinese Community Party talking points and censorship”), <https://www.whitehouse.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf>.

142. This case does not involve the kind of compelled-speech mandate that might get lesser scrutiny under *Zauderer v. Off. of Disciplinary Couns. of Supreme Ct. of Ohio*, 471 U.S. 626 (1985). *Zauderer* creates a limited exception for government-compelled speech that aims to combat misleading advertisements. It applies only to First Amendment claims involving commercial speech. See *NetChoice*, 808 F. Supp. 3d at 1239-40. And the Supreme Court has never applied the principles set forth in *Zauderer* outside the context of misleading advertising. See *NIFLA*, 585 U.S. at 768. To the contrary, the Court has consistently reaffirmed that *Zauderer* focused only on “combat[ing] the problem of inherently misleading commercial advertisements.” *Milavetz, Gallop & Milavetz, P.A. v. United States*, 559 U.S. 229, 250 (2010); see also, e.g., *Hurley*, 515 U.S. at 573, *United States v. United Foods*, 533 U.S. 405, 416 (2001).

143. *Zauderer* does not apply because SB24-205 does not compel commercial speech. Commercial speech is speech “that does no more than propose a commercial transaction.” *United Foods*, 533 U.S. at 409. That criterion is not simply a clue or guidepost; it is “*the test* for identifying commercial speech.” *City of Cincinnati v. Discovery Network, Inc.*, 507 U.S. 410, 423 (1993) (*Bd. of Trs. of State Univ. of N.Y. v. Fox*, 492 U.S. 469, 473-74 (1989)). None of the information that SB24-205 forces xAI to disclose (e.g., its practices for evaluating and mitigating model bias) has anything to do with proposing a commercial transaction, let alone addressing any potentially misleading advertising.

144. The Tenth Circuit’s multi-factor test for commercial speech, which only applies in close cases, likewise supports the conclusion that SB24-205’s disclosure provisions do not compel commercial speech. That test turns on whether the speech (1) is “concededly an advertisement”; (2) “refers to a specific product”; or (3) “is motivated by an economic interest.” *United States v.*

Wenger, 427 F.3d 840, 847 (10th Cir. 2005) (citing *Bolger v. Youngs Drug Prods. Corp.*, 463 U.S. 60, 66-67 (1983)). As an initial matter, these factors presuppose that an entity is already engaging in speech that the government then regulates—and not that the government is forcing companies to create speech. The government cannot compel private entities to speak in the first instance, and then label that speech “commercial.” Even so, the *Bolger* factors weigh against concluding that the disclosure provisions implicate commercial speech. The disclosures SB24-205 mandates are not advertisements; they are not designed to sell a product or service. They do not refer to a particular product; they describe xAI’s practices for evaluating and mitigating algorithmic bias. And the disclosures would only be produced out of legal obligation, not economic motivation.

145. Moreover, *Zauderer* does not apply because SB24-205’s compelled disclosures are neither “purely factual” nor “uncontroversial.” A disclosure is “purely factual” if it only requires “the disclosure of accurate, factual information.” *Nat’l Ass’n of Wheat Growers v. Bonta*, 85 F.4th 1263, 1276 (9th Cir. 2023). The word “purely” is critical: it excludes speech that embeds normative framing, subjective interpretation, or predictive judgment. SB24-205’s disclosures are not purely factual because they require developers to adopt and apply the statute’s definition of “algorithmic discrimination,” a normative framework shaped by Colorado’s own views about equality and fairness. Evaluating the risk of algorithmic discrimination and choosing appropriate mitigation measures is also an inherently subjective exercise. And, tellingly, the disclosures here bear little resemblance to those courts have concluded are “purely factual.” *See, e.g., Zauderer*, 471 U.S. at 652 (clarification that client may be liable for costs for attorney advertisements that promised “no recovery, no legal fees”); *Am. Meat Inst. v. U.S. Dep’t of Agric.*, 760 F.3d 18, 27 (D.C. Cir. 2014) (en banc) (country-of-origin labeling); *New York State Rest. Ass’n v. New York City Bd. of Health*,

556 F.3d 114, 131-32 (2d Cir. 2009) (calorie count); *Nat'l Elec. Mfrs. Ass'n v. Sorrell*, 272 F.3d 104, 107 (2d Cir. 2001) (instructions for safe product use and disposal).

146. SB24-205's compelled disclosures are also not "uncontroversial." A statement is "uncontroversial" under *Zauderer* "where the truth of the statement is not subject to good-faith scientific or evidentiary dispute and where the statement is not an integral part of a live, contentious political or moral debate." *Free Speech Coal., Inc. v. Paxton*, 95 F.4th 263, 281-82 (5th Cir. 2024). SB24-205's disclosures force developers to opine on the public and scholarly debate about whether using AI tools to make hiring (and other) decisions produces discriminatory results. Indeed, the disclosures implicate one of the most persistent and divisive debates in American history—what qualifies as discrimination. And xAI does not share Colorado's views—that algorithmic discrimination can be defined to exclude using AI tools to "expan[d] an applicant, customer, or participant pool to increase diversity or redress historical discrimination." § 6-1-1701(1)(b)(I)(B).

147. The sheer volume of scholarly literature on how to evaluate and mitigate algorithmic discrimination shows that SB24-205's disclosures are not "uncontroversial." Several prominent papers catalog dozens of potential fairness metrics and distinct bias types, highlighting the lack of consensus on how to measure or remediate bias. *See* Gallegos et al., *supra*, at 34-56; Guo et al., *supra*, at 15-18. Other papers underscore the difficult accuracy trade-offs that developers must weigh when considering measures designed to mitigate any bias that exists in the training datasets. *See* Tamkin et al., *supra*, at 9-11. In fact, the Association for Computing Machinery hosts an annual interdisciplinary conference focused on the social and technical challenges of algorithmic systems, including the potential for bias. In short, scholars disagree profoundly about whether AI systems produce discriminatory results, how to measure algorithmic

bias, what causes bias in outputs, and the best strategies for mitigating bias. SB24-205's compelled disclosures would force xAI to take sides in this public debate by disclosing documentation describing its own methodology for evaluating and mitigating algorithmic discrimination.

148. That these issues are also the subject of active litigation further confirms that SB24-205's compelled disclosures are not "uncontroversial" under *Zauderer*. See, e.g., *Mobley v. Workday Inc.*, 740 F. Supp. 3d 796 (N.D. Cal. 2024) (disparate impact claim under Title VII based on allegation that use of Workday's AI system produced discriminatory hiring decisions); *Huskey v. State Farm Fire & Cas. Co.*, No. 22 C 7014, 2023 WL 5848164 (N.D. Ill. Sept. 11, 2023) (Fair Housing Act claim based on allegation that use of machine-learning algorithm used to screen for potentially fraudulent claims subjected black policyholders to additional administrative hurdles and processing delays); *Louis v. SafeRent Sols., LLC*, 685 F. Supp. 3d 19 (D. Mass. 2023) (Fair Housing Act claim based on allegation that tenant-screening algorithm biased against certain applicants).

149. SB24-205's disclosure provisions trigger strict scrutiny because they are content- and viewpoint-based. And they cannot withstand strict scrutiny or any other form of heightened scrutiny. *Zauderer* does not apply because SB24-205's disclosure provisions do not compel commercial speech and, at any rate, the disclosures are neither "factual" nor "uncontroversial." SB24-205's disclosure provisions should accordingly be declared unconstitutional, and the Attorney General should be enjoined from enforcing them against xAI.

COUNT III

Declaratory Relief and Preliminary and Permanent Injunctive Relief for Extraterritorial Regulation in Violation of the Commerce Clause of the United States Constitution and the Horizontal Separation of Powers (U.S. Const. art. I, § 8, cl. 3)

150. xAI re-alleges and incorporates by reference all allegations set forth above.

151. The Constitution restricts the power of states to directly regulate conduct that takes place entirely in another state. That bedrock principle of equal sovereignty among the states is apparent in several of the Constitution’s structural protections and deeply rooted in our nation’s historical tradition. *See Nat’l Pork Producers Council v. Ross*, 598 U.S. 356, 376 n.1 (2023); *id.* at 403-04, 408-10 (Kavanaugh J., concurring in part and dissenting in part).

152. Although this principle does not erect a *per se* bar against laws that regulate conduct within one state in ways that have an “extraterritorial *effect*” in others, *Ross*, 598 U.S. at 357-58 (emphasis added), the same is not true of laws that “*directly* regulated out-of-state transactions.” *Id.* at 376 n.1. “[O]riginal and historical understandings of the Constitution’s structure and the principles of ‘sovereignty and comity’ it embraces” inform “the territorial limits of state authority under the Constitution’s horizontal separation of powers.” *Id.* at 376 & 376 n.1.

153. It is axiomatic that “all States enjoy equal sovereignty.” *Shelby Cnty. v. Holder*, 570 U.S. 529, 535 (2013); *see also PPL Mont., LLC v. Montana*, 565 U.S. 576, 591 (2012) (“the States in the Union are coequal sovereigns under the Constitution”). Indeed, “the constitutional equality of the states is essential to the harmonious operation of the scheme upon which the Republic was organized.” *Coyle v. Smith*, 221 U.S. 559, 580 (1911). When a State reaches beyond its own borders to “directly regulate[] out-of-state transactions by those with no connection to the State,” *Ross*, 598 U.S. at 376 n.1 (emphasis omitted), it invades the sovereignty and impinges on the equality of other states. Accordingly, one State is necessarily prohibited from directly regulating conduct that neither occurs nor is directed within its borders. *Cf. PennEast Pipeline Co. v. New Jersey*, 141 S.Ct. 2244, 2259 (2021) (“The plan of the Convention reflects the ‘fundamental postulates implicit in the constitutional design.’” (quoting *Alden v. Maine*, 527 U.S. 706, 729

(1999)); *State Farm Mut. Auto. Ins. Co. v. Campbell*, 538 U.S. 408, 422 (2003) (“A basic principle of federalism is that each State may make its own reasoned judgment about what conduct is permitted or proscribed within its borders, and each State alone can determine what measure of punishment, if any, to impose on a defendant who acts within its jurisdiction.”).

154. Consistent with our founding’s “special concern both with the maintenance of a national economic union unfettered by state-imposed limitations on interstate commerce and with the autonomy of the individual States within their respective spheres,” the Supreme Court has held that the Commerce Clause (U.S. Const. art. I, § 8, cl. 3) prohibits any state from “control[ling] commerce occurring wholly outside [its] boundaries.” *Healy v. Beer Inst., Inc.*, 491 U.S. 324, 335-36 (1989) (footnote omitted); *see also, e.g., Brown-Forman Distillers Corp. v. N.Y. State Liquor Auth.*, 476 U.S. 573, 579 (1986); *Edgar v. MITE Corp.*, 457 U.S. 624, 643 (1982).

155. The Court has also long interpreted the Due Process Clause to impose restrictions on a state’s ability to regulate conduct occurring wholly outside its borders. *See, e.g., Watson v. Empps. Liab. Assurance Corp.*, 348 U.S. 66, 70 (1954) (recognizing “the due process principle that a state is without power to exercise ‘extra territorial jurisdiction,’ that is, to regulate and control activities wholly beyond its boundaries”); *Home Ins. Co. v. Dick*, 281 U.S. 397, 407-08 (1930) (recognizing and applying the same principle).

156. And, of course, the Tenth Amendment provides that “powers not delegated to the United States by the Constitution, nor prohibited by it to the States, are reserved to the States respectively, or to the people,” U.S. Const. amend. X, making clear that each state retains its *own* “integrity, dignity, and residual sovereignty,” *Bond v. United States*, 564 U.S. 211, 221 (2011).

157. Despite these prohibitions, SB24-205 creates liability for xAI based purely on out-of-state conduct. The law broadly applies to any “developer”—“a person doing business in [Colorado] that develops ... an artificial intelligence system.” § 6-1-1701(7). Each of SB24-205’s provisions further require that they impact a “consumer”—*i.e.*, any “individual who is a Colorado resident.” § 6-1-1701(4). The provisions that require developers to mitigate and disclose risks concerning “algorithmic discrimination” cover “high-risk [AI] system[s],” which are systems that “make[], or is a substantial factor in making, a consequential decision” affecting “any consumer.” §§ 6-1-1701(1), (3), (4), (9)(a); 6-1-1702. Likewise, SB24-205 requires developers to disclose “to each consumer who interacts with the artificial intelligence system that the consumer is interacting with an artificial intelligence system.” § 6-1-1704(1). There is no requirement that the Colorado resident be in Colorado or that the use of the AI system or “consequential decision” occur within Colorado.

158. Thus, xAI is subject to SB24-205’s provisions merely if (1) it is a “developer” and (2) its AI system impacts a consequential decision affecting a Colorado resident. It is unimportant where that decision occurs, where the Colorado resident is located, how unforeseeable “algorithmic discrimination” was to xAI, or even whether the discrimination was intended.

159. It is not difficult to think of an example reflecting how extraordinary these provisions are in scope. Consider a healthcare company that is not incorporated or headquartered in Colorado and does not maintain offices here. That company engaged in negotiations for Grok, which was developed in California, and contracted with xAI, which is Nevada-incorporated and California-based. Under SB24-205’s provisions governing “algorithmic discrimination,” xAI would be exposed to liability under the law if a single medical decision affecting a Colorado

resident visiting one of the healthcare company’s out-of-state offices is made using Grok. § 6-1-1702.

160. Not only does xAI already have nationwide enterprise customers using Grok to make consequential decisions, *see supra* ¶ 69, xAI’s enterprise business is rapidly growing, and xAI is in the process of obtaining many additional customers who would leverage Grok to assist in making decisions in other categories covered by SB24-205, including financial and lending services, insurance, and housing. xAI’s continued expansion will certainly expose it to additional liability under SB24-205.

161. But it is not only xAI that faces this reality. None of the major enterprise AI developers are headquartered or incorporated in Colorado:

- xAI is Nevada-incorporated and California-headquartered;
- Anthropic is Delaware-incorporated and California-headquartered;
- OpenAI is Delaware-incorporated and California-headquartered; and
- Google (Gemini) is Delaware-incorporated and California-headquartered.

SB24-205’s extraterritorial scope therefore is compounded.

162. SB24-205’s provision requiring xAI to disclose “to each consumer who interacts with the artificial intelligence system that the consumer is interacting with an artificial intelligence system” is equally unconstitutional. § 6-1-1704(1). Like the provisions governing “algorithmic discrimination,” § 6-1-1702, it applies regardless of where the “consumer” (*i.e.*, Colorado resident) is, and so applies to interactions with Grok nationwide.

163. SB24-205 thus regulates transactions that are consummated entirely outside of Colorado, by entities entirely outside of Colorado, and for systems developed entirely outside of

Colorado. This violates the Constitution: indeed, imposing state-law liability on out-of-state actors for actions taken entirely out of state is the definition of unconstitutional extraterritorial state regulation. *See Ross*, 598 U.S. at 376 n.1 (“law that *directly* regulated out-of-state transactions by those with *no* connection to the State” are unconstitutional); *Edgar*, 457 U.S. at 624; *Styczinski v. Arnold*, 46 F.4th 907, 913 (8th Cir. 2022) (impermissibly extraterritorial where the law covers “a Minnesota transaction[, which] includes a transaction anywhere in the world between a bullion trader and a Minnesota resident. A bullion trader could therefore become subject to and violate Minnesota law without conducting a single transaction in Minnesota.”).

COUNT IV

Declaratory Relief and Preliminary and Permanent Injunctive Relief for the Unlawful Imposition of an Undue Burden on Interstate Commerce (*Pike* Balancing)

164. xAI re-alleges and incorporates by reference all allegations set forth above.

165. In addition to prohibiting states from regulating conduct that takes place outside of their borders, the Commerce Clause also prohibits laws where “the burden imposed on interstate commerce is clearly excessive in relation to the putative local benefits.” *Green Room LLC v. Wyoming*, 157 F.4th 1196, 1209 (10th Cir. 2025); *see also Air Transp. Ass'n of Am., Inc. v. Moss*, No. 1:23-cv-02421-DDD-KAS, 2024 WL 4369137, at *5 (D. Colo. Sept. 30, 2024).

166. In *Pike*, 397 U.S. at 142, the Supreme Court articulated the factors that a Court should consider when assessing whether a law unduly burdens interstate commerce. *Pike* balancing requires the Court to consider four factors: “(1) the nature of the putative local benefits advanced by the [statute]; (2) the burden the [statute] imposes on interstate commerce; (3) whether the burden is clearly excessive in relation to the local benefits; and (4) whether the local interests can be promoted as well with a lesser impact on interstate commerce.” *Johnson & Johnson Vision*

Care, Inc. v. Reyes, 665 F. App'x 736, 744 (10th Cir. 2016) (internal quotation marks omitted).

All four factors weigh in xAI's favor here.

167. First, the burden the statute imposes on interstate commerce is significant. As explained above, *supra* at ¶¶ 150-63, SB24-205 regulates transactions consummated entirely outside of Colorado between xAI and entities outside of Colorado for systems developed entirely outside of Colorado.

168. Not only does SB24-205 apply to transactions entirely outside of the state of Colorado, but it also applies regardless of how unforeseeable use by a Colorado resident is. In order for xAI to comply with SB24-205's "algorithmic discrimination" requirements, it would be required to alter its products *nationwide* to accommodate Colorado's requirements. Moreover, to the extent complying with SB24-205 would require retraining Grok, *supra* at ¶¶ 100-02, that will require significant economic and engineering investment.

169. Second, the nature of the putative local benefit is highly speculative, at best. SB24-205 itself contains no findings and cites no evidence that AI systems are causing discrimination—much less systemic discrimination. Additionally, federal statutes, including Title VII of the Civil Rights Act, the Fair Housing Act, the Equal Credit Opportunity Act, and the Americans with Disabilities Act, and their Colorado analogs already prohibit discriminatory outcomes in the same subject matter that SB24-205 covers.

170. Third, the burden imposed by SB24-205 is clearly excessive in relation to the local benefits. AI innovation is critical to American economic and national security. As the United States government put it, AI "ha[s] the potential to reshape the global balance of power" and "it is a national security imperative for the United States to achieve and maintain unquestioned and

unchallenged global technological dominance.”⁷⁰ “United States AI companies must be free to innovate without cumbersome regulation,” “excessive State regulation thwarts this imperative,” “State-by-State regulation by definition creates a patchwork of 50 different regulatory regimes that makes compliance more challenging,” and “State laws sometimes impermissibly regulate beyond State borders, impinging on interstate commerce.”⁷¹ The White House more recently explained that “winning the AI race” would “usher in a new era of human flourishing, economic competitiveness, and national security for the American people,” “AI cannot become a vehicle for government to dictate right and wrong-think,” and “[a] patchwork of conflicting state laws would undermine American innovation and our ability to lead in the global AI race.”⁷² AI systems function similarly to the internet, rail, or highway insofar as they “require[] a cohesive national scheme of regulation so that users are reasonably able to determine their obligations.” *ACLU v. Johnson*, 194 F.3d 1149, 1162 (10th Cir. 1999). These concerns apply on all fours to SB24-205, which imposes broad and costly requirements on AI developers operating anywhere nationwide. Even the Attorney General himself has expressed concern about the impact SB24-205 might have, going as far as saying the bill “is really problematic.”⁷³

⁷⁰ <https://www.whitehouse.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf>.

⁷¹ <https://www.whitehouse.gov/presidential-actions/2025/12/eliminating-state-law-obstruction-of-national-artificial-intelligence-policy/>.

⁷² <https://www.whitehouse.gov/articles/2026/03/president-donald-j-trump-unveils-national-ai-legislative-framework/>.

⁷³ <https://broadbandbreakfast.com/state-ag-warns-colorado-ai-bill-could-drive-innovation-out-of-state/>; see also <https://newspack-coloradosun.s3.amazonaws.com/wp-content/uploads/2024/06/FINAL-DRAFT-AI-Statement-6-12-24-JP-PW-and-RR-Sig.pdf>.

171. Fourth, the putative local interests served by SB24-205 can be promoted through lesser impacts on interstate commerce. Colorado already has comprehensive discrimination laws. For example, the Colorado Anti-Discrimination Act prohibits discriminatory or unfair employment and housing practices. SB24-205 does not impair compliance with that or other anti-discrimination laws and even expressly states that nothing in it “restricts a developer’s ... ability to: (a) [c]omply with federal, state, or municipal laws, ordinances, or regulations.” § 6-1-1705(1)(a).

172. SB24-205’s requirement that AI developers prevent “algorithmic discrimination” might also have unintended effects. AI developers may incidentally overcorrect, resulting in the very discrimination SB24-205 purports to address. *See Students for Fair Admissions, Inc. v. President & Fellows of Harvard Coll.*, 600 U.S. 181, 223 (2023) (“*SFFA*”) (“Outright racial balancing is ‘patently unconstitutional’” because equal protection “command[s] that the Government must treat citizens as individuals, not as simply components of a racial, religious, sexual or national class.” (cleaned up)).

173. Further, such model-rebalancing efforts have no “logical end point.” *Id.* at 221. It is not clear at what point such efforts would satisfy compliance such that the AI developer has used “reasonable care to protect consumers from any known or reasonably foreseeable risks of algorithmic discrimination.” § 6-1-1702(1).

174. SB24-205 imposes an enormous burden on interstate commerce—regulating entirely out-of-state transactions and threatening to undermine American AI innovation—relative to its local benefit—speculative concerns about “algorithmic discrimination” that lack any legislative findings or evidence.

COUNT V

**Declaratory Relief and Preliminary and Permanent Injunctive Relief for Vagueness in
Violation of the Due Process Clause of the Fourteenth Amendment to the
United States Constitution (42 U.S.C. § 1983, 28 U.S.C. § 2201(a))**

175. xAI re-alleges and incorporates by reference all allegations set forth above.

176. “A fundamental principle in our legal system is that laws which regulate persons or entities must give fair notice of conduct that is forbidden or required.” *FCC v. Fox TV Stations*, 567 U.S. 239, 253 (2012). A law is unconstitutionally vague if it “fails to provide a person of ordinary intelligence fair notice of what is prohibited, or is so standardless that it authorizes or encourages seriously discriminatory enforcement.” *United States v. Williams*, 553 U.S. 285, 304 (2008). “When speech is involved, rigorous adherence to those requirements is necessary to ensure that ambiguity does not chill protected speech.” *Fox*, 567 U.S. at 253-54. After all, vague laws risk chilling speech by forcing speakers “to ‘steer far wider of the unlawful zone’” than they would “if the boundaries of the forbidden areas were clearly marked.” *Baggett v. Bullitt*, 377 U.S. 360, 372 (1964). For that reason, laws touching on speech must themselves speak “with narrow specificity.” *NAACP v. Button*, 371 U.S. 415, 433 (1963).

177. SB24-205 is not written “with narrow specificity,” and it fails to provide fair notice to a person of ordinary intelligence as to what it requires.

178. To begin, it fails to adequately articulate what “algorithmic discrimination” is. SB24-205 defines “algorithmic discrimination” as “an unlawful differential treatment or impact that disfavors an individual or group of individuals on the basis of their actual or perceived age, color, disability, ethnicity, genetic information, limited proficiency in the English language, national origin, race, religion, reproductive health, sex, veteran status, or other classification protected under the laws of this state or federal law.” § 6-1-1701(1)(a).

179. But serious questions are left unanswered. What constitutes a “perceived attribute”? “Perceived” by whom? Is it the perception of the individual who is allegedly being discriminated against or the AI system allegedly resulting in a disparate impact? If the former, how sincerely held does that perception need to be? How do “unlawful differential treatment” or “impact” differ? And how do those tests differ from the traditional legal frameworks for disparate treatment?

180. The vagueness of SB24-205’s “algorithmic discrimination” is further compounded by its carve-outs. § 6-1-1701(1)(b). SB24-205 provides that “‘Algorithmic discrimination’ does not include: . . . (B) [e]xpanding an applicant, customer, or participant pool to increase diversity or redress historical discrimination.” *Id.* What constitutes “historical discrimination”? How far back in time must one go? Does it only include the historical treatment of African Americans arising from slavery, or does it include discrimination against the Irish and Italians in the 19th century?⁷⁴ *Nuziard v. Minority Bus. Dev. Agency*, 721 F. Supp. 3d 431, 509 (N.D. Tex. 2024) (“In the mid-19th Century, Irish immigrants faced a common roadblock when seeking jobs after their trans-Atlantic voyage: motivated by rising populism and anti-Irish animus, many employers hung signs explaining ‘Irish Need Not Apply.’”). Likewise, what is permissible “diversity”? § 6-1-1701 defines algorithmic discrimination to include virtually every actual or perceived protected trait; which of those same traits are the appropriate subjects of State-approved “diversity” considerations?

181. Next, take SB24-205’s definition of a “high-risk artificial intelligence system,” which “means any artificial intelligence system that, when deployed, makes, or is a substantial

⁷⁴ <https://www.nytimes.com/interactive/2019/10/12/opinion/columbus-day-italian-american-racism.html>.

factor in making, a consequential decision.” § 6-1-1701(9)(a). When does an AI system become a “substantial factor”? According to SB24-205, “substantial factor” “includes *any* use of an artificial intelligence system to generate *any* content, decision, prediction, or recommendation concerning a consumer that is used as *a* basis to make a consequential decision concerning the consumer.” § 6-1-1701(11)(b) (emphasis added). That definition—“any use” of an AI system to generate “any content ... that is used as a basis to make a [] decision”—is so amorphous it encompasses virtually all uses of AI systems.

182. Now, take xAI’s obligation to prevent “algorithmic discrimination.” § 6-1-1702. SB24-205 requires xAI to use “reasonable care to protect consumers from any known or reasonably foreseeable risks of algorithmic discrimination.” “[R]easonable care” and “reasonably foreseeable” are undefined and uncertain.

183. SB24-205’s disclosure provisions fare no better. For example, SB24-205 requires developers to disclose “high-level summaries of the type of data used to train the high-risk artificial intelligence system” as well as a “general statement describing the reasonably foreseeable uses and known harmful or inappropriate uses of the high-risk artificial intelligence system.” §§ 6-1-1702(2)(b)(I), (2)(a). How “high-level” and how “general” do disclosures need to be to comply? Likewise, what constitutes “harmful or inappropriate use,” which is, again, undefined?

184. Similarly, the statute requires disclosure to the Attorney General of “known or reasonably foreseeable risks” within ninety days of discovering an AI system caused algorithmic discrimination, was “reasonably likely” to have caused algorithmic discrimination, or of receiving a “credible report that the high-risk artificial intelligence system has been deployed and has caused algorithmic discrimination.” § 6-1-1702(5)(b). When does algorithmic discrimination become

“reasonably likely”? And what constitutes a “credible report” such that an AI developer is required to comply with the disclosure provision? Would a single consumer report satisfy either trigger? If not, how many consumer reports must an AI developer receive before it takes action and what must those reports contain such that they are sufficiently “credible.” SB24-205 is long on questions and short on answers.

185. What’s worse, SB24-205 gives the Attorney General a massive amount of interpretive and rulemaking power. SB24-205 gives the Colorado Attorney General rulemaking authority to implement the law and has exclusive authority to enforce it. §§ 6-1-1706(1)(2), 6-1-1707. The Attorney General alone thus has authority to decide what actions violate SB24-205 and its nebulous provisions and determine who to enforce it against.

186. The statute’s use of vague, relative and uncertain terms gives developers no fair notice of what mitigation or disclosures are required, and hands the Colorado Attorney General virtually unfettered discretion to enforce the law arbitrarily against his political opponents. SB24-205’s terms are so malleable that it authorizes “arbitrary [or] discriminatory enforcement.” *Grayned v. City of Rockford*, 408 U.S. 104, 108 (1972).

COUNT VI

Declaratory Relief and Preliminary and Permanent Injunctive Relief for Violation of the Equal Protection Clause of the Fourteenth Amendment to the United States Constitution (42 U.S.C. § 1983, 28 U.S.C. § 2201(a))

187. xAI re-alleges and incorporates by reference all allegations set forth above.

188. “Proposed by Congress and ratified by the States in the wake of the Civil War, the Fourteenth Amendment provides that no State shall ‘deny to any person . . . the equal protection of the laws.’” “Proponents of the Equal Protection Clause described its ‘foundation[al] principle’ as ‘not permit[ing] any distinctions of law based on race or color.’” “Any ‘law which operates

upon one man,’ they maintained, should ‘operate equally upon all.’ Accordingly, as th[e] [Supreme] Court’s early decisions interpreting the Equal Protection Clause explained, the Fourteenth Amendment guaranteed ‘that the law in the States shall be the same for the black as for the white; that all persons, whether colored or white, shall stand equal before the laws of the States.’” *SFFA*, 600 U.S. at 183 (citations omitted). “Eliminating racial discrimination means eliminating all of it.” *Id.* at 183-84.

189. “Given its drafters’ desire to eschew race-based distinctions, even benign discrimination cannot be squared with the Equal Protection Clause—it can only be an ‘exception.’” *Nuziard*, 721 F. Supp. 3d at 477-78 (citation omitted). “Any exceptions to the Equal Protection Clause’s guarantee must survive a daunting two-step examination known as ‘strict scrutiny,’ which asks first whether the racial classification is used to ‘further compelling governmental interests,’ and second whether the government’s use of race is ‘narrowly tailored,’ *i.e.*, ‘necessary,’ to achieve that interest.” *SFFA*, 600 U.S. at 184 (citations omitted). SB24-205’s requirement to mitigate “algorithmic discrimination” flunks that test.

190. SB24-205 requires xAI to mitigate against and make disclosures concerning “algorithmic discrimination.” §§ 6-1-1701(1), 1702. According to Colorado, however, not all algorithmic discrimination is impermissible “algorithmic discrimination.” While xAI has a general duty to mitigate against “algorithmic discrimination,” “[a]lgorithmic discrimination’ does not include: (I) [t]he offer, license, or use of a high-risk artificial intelligence system by a developer or deployer for the sole purpose of ... (B) [e]xpanding an applicant, customer, or participant pool to increase diversity or redress historical discrimination.” § 6-1-1701(1)(b)(I)(B). In other words, if xAI engages in balancing to “increase diversity” or “redress historical discrimination,” including

on the basis of “race,” it does not engage in “algorithmic discrimination” within the meaning of SB24-205 at all. *Id.*; *see also* § 6-1-1701(1).

191. Accordingly, if an AI developer alters its model to provide preferencing mandated by Colorado to increase diversity or redress historical discrimination, it is effectively exempt from SB24-205’s mitigation requirements. That is because if an AI developer alters its AI system to increase “diversity” or “redress historical discrimination” then any disparate treatment that occurs as a result does not constitute “algorithmic discrimination.” Therefore, the AI developer need not use “reasonable care to protect consumers from any known or reasonably foreseeable risks of algorithmic discrimination” because no such “algorithmic discrimination” could occur. § 6-1-1702.

192. SB24-205 thus presents xAI with a choice. xAI can either continue to develop Grok as a maximally truth-seeking AI system that provides objective and unbiased outputs; expend significant resources on compliance, documentation, and risk mitigation; and face full enforcement risk under SB24-205 for any disparate impact. Or it can build an AI with the sole purpose of increasing diversity or redressing historical discrimination, and receive a safe harbor from potential liability under SB24-205. That is an impermissible choice. Affording benefits to some, but not others, based on the treatment of protected characteristics, including race, is precisely the type of “[d]istinction[] between citizens solely because of their ancestry [that] are by their very nature odious to a free people whose institutions are founded upon the doctrine of equality.” *SFFA*, 600 U.S. at 184 (citations omitted).

193. SB24-205 includes no justification, much less a compelling justification, for its codified discrimination. Again, it includes no statement of purpose. SB24-205 does not even define what “diversity” is or what “historical discrimination” it seeks to redress.

194. At best, SB24-205’s discrimination provision seems to be designed to embody Colorado’s subjective vision of equity and remedy amorphous discrimination across a wide range of services and opportunities, including education, employment, financial or lending services, health-care services, housing, essential government services, insurance and legal services. § 6-1-1701(3). But “an effort to alleviate the effects of societal discrimination is not a compelling interest.” *Shaw v. Hunt*, 517 U.S. 899, 909 (1996). If it was, the Equal Protection Clause “would be lost in a mosaic of shifting preferences based on inherently unmeasurable claims of past wrongs.” *City of Richmond v. J.A. Croson Co.*, 488 U.S. 469, 505-06 (1989). “The mere recitation of a benign or compensatory purpose for the use of a racial classification would essentially entitle the States to exercise the full power of Congress under § 5 of the Fourteenth Amendment and insulate any racial classification from judicial scrutiny under § 1. We believe that such a result would be contrary to the intentions of the Framers of the Fourteenth Amendment, who desired to place clear limits on the States’ use of race as a criterion for legislative action, and to have the federal courts enforce those limitations.” *Id.* at 490-91 (citations omitted).

195. To the extent that Colorado believes SB24-205 addresses specific instances of discrimination, it must provide facts to support that conclusion. “First, it must identify the past or present discrimination ‘with some specificity.’ Second, it must also demonstrate that a ‘strong basis in evidence’ supports its conclusion that remedial action is necessary.” *Concrete Works of Colo. v. City & Cty. of Denver*, 321 F.3d 950, 958 (10th Cir. 2003) (citations omitted); *see also*

Nuziard, 721 F. Supp. 3d at 480. Colorado “cannot point to general social ills and call it a day. Rather, it must identify the ‘who, what, when, where, why, and how’ of relevant discrimination. *Nuziard*, 721 F. Supp. 3d at 480 (citing *Croson*, 488 U.S. at 492). Yet SB24-205 includes no legislative findings or evidence whatsoever, much less any that would prove the existence of specific harm that the state has a compelling interest in redressing.

196. SB24-205 is also not narrowly tailored. In fact, the provision excepting permissible “algorithmic discrimination” applies to any protected characteristic so long as their inclusion increases “diversity” or redresses “historical discrimination.” § 6-1-1701(1)(b)(I)(B). It also applies to virtually any decision that yields disparate results because SB24-205 covers nearly every subject matter—education, employment, finance, government services, healthcare, housing, insurance, and legal services. § 6-1-1701(3). That is woefully overinclusive. SB24-205 encodes State preferencing for virtually *any* protected characteristic across virtually *every* subject matter. Take employment for example: “If a [preference] was ‘narrowly tailored’ to compensate black contractors for past discrimination, one may legitimately ask why they are forced to share this ‘remedial relief’ with an Aleut citizen who moves to [Boulder] tomorrow?” *Croson*, 488 U.S. at 506.

197. SB24-205 is overinclusive for another reason: its carveout for “diversity” and “historical discrimination” paints in broad strokes and encompasses members with protected traits who may personally experience no discrimination or disparate impact. As Governor Polis explained when expressing concern over the bill, “[l]aws that seek to prevent discrimination

generally focus on prohibiting intentional discriminatory conduct,” but SB24-205 “deviates from that practice by regulating the results of AI system use, regardless of intent.”⁷⁵

198. SB24-205 is also underinclusive because the definition of “algorithmic discrimination” does not include protections for individuals or groups that may actually face disparate impact in the covered subjects (housing, finance, legal services, etc.), including the economically disadvantaged. As one Court succinctly put it, “Oprah Winfrey is presumptively disadvantaged, while ... even more disadvantaged Americans are not.” *Nuziard*, 721 F. Supp. 3d at 491.

199. SB24-205 also lacks a “logical end point.” *SFFA*, 600 U.S. at 221. SB24-205 includes no guidance indicating when “increas[ing] diversity” or “redress[ing] historical discrimination” will be sufficiently addressed by changes in xAI’s model such that further inclusion of protected groups would actually *decrease* “diversity” or *affirmatively cause* “historical discrimination” of excluded majority groups. SB24-205’s extraordinarily broad affirmative discrimination provisions make it “unclear how a court,” or xAI for that matter “is supposed to determine if or when such goals would be adequately met.” *Id.*

200. Lastly, as explained above, less restrictive alternatives exist. State and federal law—for example, the Civil Rights Act of 1964, Fair Housing Act, Equal Credit Opportunity Act, and Colorado Anti-Discrimination Act—already prohibit intentional discrimination in employment, housing, education, finance, and other decisions within SB24-205’s scope.

⁷⁵ Jared Polis, Statement on Signing SB24-205 (May 17, 2024), <https://drive.google.com/file/d/1i2cA3IG93VViNbZXu9LPgbTrZGqhyRgM/view>.

201. SB24-205’s carveout for permissible “algorithmic discrimination” is as amorphous as it is unconstitutional, and forces xAI to choose between adopting Colorado’s state-ordained preferencing in its AI model or continuing to operate its maximally-truth seeking AI and expose itself to liability, including civil penalties.

PRAYER FOR RELIEF

For the foregoing reasons, xAI respectfully requests that the Court:

A. Enter a judgment declaring, pursuant to 28 U.S.C. §§ 2201(a) and 2202, that SB24-205 unconstitutionally compels xAI’s speech in violation of the First Amendment of the U.S. Constitution;

B. Enter a judgment declaring, pursuant to 28 U.S.C. §§ 2201(a) and 2202, that SB24-205 constitutes an impermissible extraterritorial regulation;

C. Enter a judgment declaring, pursuant to 28 U.S.C. §§ 2201(a) and 2202, that SB24-205 imposes an undue burden on interstate commerce;

D. Enter a judgment declaring, pursuant to 28 U.S.C. §§ 2201(a) and 2202, that SB24-205 is unconstitutionally vague in violation of the Due Process Clause of the U.S. Constitution;

E. Enter a judgment declaring, pursuant to 28 U.S.C. §§ 2201(a) and 2202, that SB24-205 violates the Equal Protection Clause of the Fourteenth Amendment of the U.S. Constitution;

F. Issue an order permanently enjoining Attorney General Weiser, as well as all officers, agents, and employees subject to his supervision, direction, or control, from enforcing the provisions of SB24-205 against xAI;

G. Award xAI all costs and attorney’s fees to which it is entitled pursuant to 42 U.S.C. § 1988 and other applicable law; and

H. Grant xAI such other and further relief as the Court deems just and proper.

Dated: April 9, 2026.

Respectfully submitted,

s/ Frederick R. Yarger

Frederick R. Yarger
William D. Hauptman
Wheeler Trigg O'Donnell LLP
370 Seventeenth Street, Suite 4500
Denver, CO 80202
Telephone: 303.244.1800
Facsimile: 303.244.1879
Email: yarger@wtotrial.com
hauptman@wtotrial.com

James Burnham
X.AI LLC
601 13th Street NW 12th Floor
Washington, DC 20005
Email: jburnham@x.ai

Adam Mehes
Argirios J. Nickas
X.AI LLC
249 W 17th Street
New York, NY 10011
Email: amehes@x.ai
anickas@x.ai

Reid Coleman
X.AI LLC
865 FM 1209, Building 2
Bastrop, TX 78602
Email: rcoleman@x.ai

Attorneys for Plaintiff X.AI LLC