

A BIPARTISAN POSITION PAPER AND LEGISLATIVE FRAMEWORK

Markets Need Referees:

An SEC-Model Commission for Artificial Intelligence

Proposing the

**American Artificial Intelligence Leadership and Accountability Act of
2026**

establishing the Artificial Intelligence Oversight Commission

Published jointly by

THE ALIGN AI FOUNDATION

and

THE GIE FOUNDATION

A working draft for congressional consideration

July 2026

Contents

Foreword

This paper is issued jointly by two organizations that arrived at the same conclusion from opposite ends of the problem.

The **Align AI Foundation** funds the research any credible oversight regime must eventually rely on: interpretability work that makes model behavior legible, secure environments for evaluating systems before they reach the public, and the study of deception and concealed capability in advanced models. That research has produced an uncomfortable finding, and it is the premise of this paper — alignment failures are not a future hypothesis. They are documented, measurable, and already reaching people. What research cannot supply is an institution with the standing to receive what it learns and the authority to act on it.

The **Geopolitical Insight and Education Foundation** studies the drivers of instability in democratic societies — technological, environmental, and geopolitical — and how institutions built for a slower century absorb them. Artificial intelligence has become the clearest instance of the pattern GIEF was founded to examine: a technology advancing faster than the governing capacity of the institutions responsible for it, and governed in the interim by litigation, executive improvisation, and forty state legislatures working at cross purposes.

Neither half of the problem is solvable alone. Technical safety knowledge without institutional capacity is inert — findings that no one is obligated to receive. Institutional design without technical grounding is theater — process that cannot distinguish a real safety case from a well-drafted one. What follows is our attempt to write the intersection: a framework grounded in what evaluation science can actually verify today, housed in the most durable regulatory design in American law.

We are independent, and the proposal reflects it. Align AI accepts no corporate revenue and holds no position in the industry it studies. GIEF is not funded by, or beholden to, any government, political party, or interest. That independence is why this document proposes measures much of the industry will resist, and stops short of measures many safety advocates would prefer. We have followed the evidence where it leads, and it leads to disclosure, examination, and accountable process — not government approval of products, and not another decade of voluntary commitments.

The framework is offered as a working draft, in the plainest sense of the term. Every threshold, clock, and definition in it is open to challenge, and we will publish revisions as the evidence changes. We ask only that the alternative be judged on the same terms: the status quo is not the absence of a decision about how artificial intelligence is governed. It is a decision to let the least accountable actors keep making it.

Gary Kucher

Founder

The Align AI Foundation

Bennett Iorio

Founder and Director

Geopolitical Insight and Education Foundation

Executive Summary

The United States is regulating the most consequential technology of the century through a patchwork of state statutes, executive orders, consent decrees, and voluntary commitments — and everyone involved knows it. Forty states have enacted artificial intelligence laws; the leading bipartisan proposal in Congress would trade federal preemption for a safety framework administered by no dedicated agency; and the executive branch, having twice watched Congress reject preemption without substance (including by a 99–1 Senate vote), has resorted to voluntary pre-release arrangements with frontier developers. Meanwhile, **80 percent of Democratic voters and 84 percent of Republican voters support creating a new federal agency to oversee artificial intelligence** — a level of cross-party agreement almost no other live policy question can claim.

This paper proposes that Congress do for artificial intelligence what it did for securities markets in 1934: create an expert, bipartisan, independent commission — the **Artificial Intelligence Oversight Commission (AIOC)** — through the **American Artificial Intelligence Leadership and Accountability Act of 2026**. The Commission would be modeled structurally and philosophically on the Securities and Exchange Commission: five members, no more than three from one party, staggered terms, funded substantially by industry fees, and built on the SEC's core regulatory insight — *mandated disclosure and accountable process, not government product approval*. The SEC does not decide which stocks are good investments. The AIOC would not decide which AI models are good products. It would ensure that the developers of the most capable systems tell the truth, in a standard format, on a fixed schedule, with real penalties for deception — and that catastrophic-risk claims are examined by someone other than the companies making them.

The framework in brief

- **Tiered by capability and risk.** The overwhelming majority of AI developers, deployers, startups, researchers, and open-source projects face no obligations at all. Registration and annual disclosure apply only to large covered developers; the full regime — safety-case filings, incident reporting, examinations — applies only to frontier developers above high compute and revenue thresholds, adjusted biennially.
- **Disclosure and registration at the core.** Covered developers file standardized annual disclosures (a Form 10-K analog); frontier developers additionally file a pre-deployment Frontier Model Safety Case (an S-1 analog) that becomes effective automatically unless the Commission acts within a fixed 90-day clock — a default-to-market design borrowed directly from securities registration.
- **Incident reporting.** Critical safety incidents must be reported within 72 hours (an 8-K analog), with a good-faith safe harbor, replacing today's scattered obligations to California, New York, and foreign regulators with one federal filing.
- **Examination, audit, and enforcement.** Third-party audit standards, examination authority over registered developers, civil penalties, cease-and-desist power, and a Dodd-Frank-style whistleblower program. No new private right of action.

- **A supervised self-regulatory organization.** Authority to register a FINRA-style SRO — the structure reportedly favored by senior administration officials and proposed publicly by industry leaders in July 2026 — but under statutory Commission oversight rather than as a substitute for it.
- **Partial preemption.** Developers in good standing receive uniform national rules for how models are built and disclosed, preempting conflicting state development-level mandates — while states fully retain consumer protection, civil rights, tort, child-safety, and criminal authority. Substance for preemption: the trade Congress has signaled it will accept, and the bare preemption it has twice rejected.
- **Innovation guaranteed in statute.** A regulatory sandbox, small-entity exemptions, capability-based (not openness-based) treatment of open-weight models, a statutory Office of Innovation, biennial GAO review, and a seven-year reauthorization requirement.

The moment is unusually favorable. Congressional Republicans and Democrats have each rejected the other's least-acceptable outcome — unregulated preemption and an unpreempted patchwork, respectively. The administration's own June 2026 executive order establishing voluntary pre-release model access concedes the oversight gap while demonstrating that voluntariness is the only tool the executive branch has. Industry itself is now proposing SEC-adjacent structures. What no existing proposal supplies is the institution: a durable, expert, bipartisan referee. This paper supplies the blueprint.

I. The Case for Action

Every previous general-purpose technology that reshaped American markets eventually acquired an institutional referee: railroads and the ICC, securities and the SEC, aviation and the FAA, telecommunications and the FCC, pharmaceuticals and the FDA. In each case, Congress acted not at the technology's birth but at the moment its market grew too large, too fast-moving, and too consequential for common-law remedies and state-by-state rules to govern — and in each case, the availability of trusted rules *accelerated* the industry's growth rather than arresting it. Investors returned to U.S. capital markets after 1934 because disclosure and enforcement made them the most trusted in the world. That trust premium is now America's to win or lose in artificial intelligence.

Artificial intelligence has reached that moment, and four facts define it.

First, the harms are no longer hypothetical. In January 2026, Character.AI and Google settled a set of wrongful-death and injury suits brought by families in four states, including the family of a 14-year-old who died by suicide after prolonged interaction with a companion chatbot; seven additional suits filed in November 2025 allege that a leading chatbot acted, in the plaintiffs' words, as a "suicide coach." In the winter of 2025–26, an image-generation feature on a major platform produced sexualized imagery of real, identifiable people at mass scale — advocacy analyses estimated millions of images in a matter of weeks, including apparent child sexual abuse material — triggering investigations on three continents. Congress has responded piecemeal and admirably: the TAKE IT DOWN Act became law in May 2025, and the GUARD Act protecting minors from AI companions advanced out of the Senate Judiciary Committee unanimously in April 2026. But incident-by-incident lawmaking is a symptom of missing institutional capacity, not a substitute for it.

Second, the state patchwork is real, and it is unstable in both directions. Forty states enacted at least one AI law in 2025; by March 2026, forty-five states had introduced over 1,500 AI bills. Four states — California, Colorado, New York, and Texas — have enacted four materially different frontier-model or comprehensive AI regimes with four different scopes and effective dates. Colorado's landmark 2024 law was delayed once and then substantially rewritten in May 2026 before it ever took effect. Industry cannot plan against rules that differ by state line and change by special session; safety advocates cannot rely on protections that evaporate under pressure. Both camps are right, and both are pointing at the same missing federal institution.

Third, Congress has already defined the acceptable deal — twice. The Senate stripped a ten-year state-AI moratorium from the 2025 reconciliation bill by a vote of 99–1, and a second preemption attempt failed in the FY2026 defense authorization. The message was not that Congress opposes national uniformity; it is that Congress will not grant preemption *in exchange for nothing*. The leading bipartisan proposal now circulating — the Obernolte–Trahan discussion draft of June 2026 — accepts exactly this logic, pairing a federal safety framework for large developers with a three-year freeze on state development-level laws. What that draft conspicuously lacks is an institution to administer the bargain. A framework without an agency is a statute without a referee.

Fourth, the demand for a referee now comes from every direction. The administration's June 2, 2026 executive order created a classified benchmarking process for frontier models and a voluntary framework under which developers may grant the government up to 30 days of pre-release model access — while expressly disclaiming any mandatory licensing or preclearance authority, because no such authority exists in law. Senior administration officials are reported to favor a FINRA-style oversight body connected to the SEC. The chief executive of a leading AI laboratory publicly proposed a supervised industry standards body for frontier models in July 2026. And the public is far ahead of Washington: 80 percent of Democrats and 84 percent of Republicans support a new federal AI oversight agency, and more than eight in ten voters of both parties say companies should not deploy smarter-than-human systems until they can demonstrate control. When the executive branch, industry leadership, and four-fifths of both parties' voters converge on the same institutional gap, the remaining question is not whether to build the institution but how to design it well.

This paper answers that question with the most successful template in the American regulatory tradition: the Securities and Exchange Commission.

II. The Current Landscape

A. Congress: frameworks without an institution

The 119th Congress has produced serious bipartisan work product, none of it yet law, and none of it institutional. The **Great American Artificial Intelligence Act of 2026** (Obernolte–Trahan discussion draft, June 4, 2026) is the most comprehensive: 269 pages establishing the first federal safety framework for large AI developers in exchange for a three-year moratorium on state laws governing how models are built, with states retaining authority over use. It drew early support from the Speaker and immediate opposition from labor and civil-society groups — but its obligations would be administered through existing departments, with no independent commission, no examination corps, and no dedicated technical staff. The **GUARD Act** (Hawley–Blumenthal, S. 3062) advanced from Judiciary unanimously in April 2026, addressing AI companions and minors. The **SANDBOX Act** (Cruz, S. 2750) would let AI firms obtain renewable waivers from federal regulations — a useful innovation mechanism in search of a regulator to administer it. Additional pending measures address foundation-model transparency (H.R. 8094), federal agency use of the NIST AI Risk Management Framework (H.R. 8819), and standards and evaluation (S. 3952). Each bill presumes evaluative and administrative capacity that no federal body currently possesses.

B. The executive branch: conceding the gap

The administration's July 2025 *AI Action Plan* set a deregulatory, innovation-first course, and the renamed Center for AI Standards and Innovation (CAISI) at Commerce conducts national-security evaluations of frontier models — as an evaluator, without rulemaking or enforcement power. In December 2025, following Congress's refusal to preempt by statute, an executive order directed agencies toward a uniform national framework and established a Department of Justice task force

to challenge state AI laws in court — preemption by litigation, the weakest and most contested form. Most tellingly, the **June 2, 2026 executive order** on advanced AI created a classified process to identify covered frontier models and a voluntary framework for up to 30 days of pre-release government access, while expressly stating that it shall not be construed to authorize mandatory licensing, preclearance, or permitting. The executive branch has, in effect, published its own gap analysis: it wants pre-deployment visibility into frontier systems, believes it lacks statutory authority to require it, and is improvising with volunteers. Voluntary access endures exactly as long as every frontier developer finds it convenient.

C. The states: forty laboratories, no control group

The state record since 2024 demonstrates both genuine policy energy and structural instability:

- **California** — the Transparency in Frontier Artificial Intelligence Act (SB 53), effective January 1, 2026, requires large frontier developers (training runs above 10^{26} FLOP and revenue above \$500 million) to publish safety frameworks and report critical incidents to state emergency services, with whistleblower protections — alongside training-data disclosure (AB 213) and companion-chatbot rules (SB 243).
- **Colorado** — the first comprehensive state AI act (2024) was delayed to mid-2026, then rewritten in May 2026 into a narrower disclosure regime effective 2027, eliminating its duty of care and impact-assessment mandates before they ever operated.
- **New York** — the RAISE Act (signed December 2025, amended March 2026, effective January 2027) imposes frontier-developer safety-plan and incident-reporting duties, enacted in direct defiance of the federal preemption executive order.
- **Texas** — TRAIGA (effective January 1, 2026) takes an intent-based prohibitory approach with a regulatory sandbox and attorney-general enforcement.

Four frontier regimes, four thresholds, four effective dates — two of which moved after enactment — plus an executive-branch litigation campaign against all of them. This is the environment in which American developers currently make multi-billion-dollar training decisions. A national developer today owes overlapping incident reports to Sacramento and Albany on different definitions and clocks; under the framework proposed here, one federal filing would satisfy all of them.

D. The international context

The European Union's AI Act is in force, with general-purpose model obligations applicable since August 2025 and chatbot-disclosure and synthetic-media labeling duties applying August 2, 2026 — though the EU's June 2026 "Digital Omnibus" postponed most high-risk use-case obligations to late 2027, a reminder that even Brussels found its own prescriptive machinery too heavy to implement on schedule. China mandates national provenance labeling of AI content and enforces it. The United Kingdom continues a sectoral, non-statutory approach. The strategic point for Congress is not to copy any of these. It is that the world's other major AI jurisdictions are building *institutional capacity to see inside frontier systems*, and the United States — home to the leading developers — has none. An American disclosure-based regime, run by an expert commission with

real independence, would set the global benchmark before Brussels' model becomes the default export standard, exactly as U.S. securities disclosure became the world's template after 1934.

III. Why the SEC Model

A. 1933 and 1934: the sequence we are living through again

The SEC analogy is precise in a way that is easy to miss. After the 1929 crash, Congress did not begin with an agency. It began with the **Securities Act of 1933** — a pure disclosure statute, administered by the already-existing Federal Trade Commission, requiring issuers to register offerings and tell the truth about them. Only a year later, having seen that disclosure rules without a dedicated institution could not police an industry this large, did Congress pass the **Securities Exchange Act of 1934** and create the Commission itself: five members, no more than three from one party, staggered terms, supervising both issuers and industry self-regulatory bodies.

American AI policy is currently at its 1933 moment. California's SB 53 and New York's RAISE Act are disclosure statutes administered by pre-existing, non-specialist offices. The Obernolte–Trahan draft would federalize the disclosure layer. The administration's executive orders improvise the supervisory layer through voluntary arrangements. Every element of 1934 is being reinvented ad hoc — the disclosure forms, the incident reports, the industry self-regulatory proposals, even the enforcement litigation — except the Commission that makes the system coherent. This paper proposes completing the sequence deliberately rather than accidentally.

B. The core insight: disclosure, not approval

The SEC's philosophical settlement is the reason it commands durable support across administrations of both parties: *the government does not judge the merits of the product; it polices the honesty and completeness of what the seller says about it, and it acts against fraud and defined categories of dangerous conduct.* Applied to AI, this settlement resolves the objection that has stalled every licensing proposal since 2023 — that a federal agency would become a product-approval gatekeeper, an "FDA for AI" deciding which models Americans may use. Under the framework proposed here, the Commission never approves a model. It ensures that frontier developers' safety claims are standardized, evidenced, examined, and honest; that dangerous-capability findings and serious incidents surface to someone with authority to act; and that deception about safety is punished the way deception about earnings is punished. The default is deployment. The burden of stopping a release rests on the agency, on defined findings, on a statutory clock, subject to expedited judicial review — the exact mechanics of a securities registration statement, which becomes effective by operation of law unless the SEC acts.

C. Mapping the machinery

The table below maps each proven SEC mechanism to its AI analog in the proposed framework.

SEC mechanism	Securities function	AI analog (proposed)
Issuer registration	Public companies register and are known to the regulator	Covered AI developers register with the Commission (Title II)
Form 10-K / periodic reporting	Standardized annual disclosure of material facts	Annual AI disclosure report: capabilities, safety practices, evaluations, compute (Title II)
S-1 registration statement	Pre-offering disclosure; effective unless SEC acts; stop-order power	Frontier Model Safety Case filed pre-deployment; deemed effective in 90 days absent Commission action (Title III)
Form 8-K	Prompt disclosure of material events	72-hour critical safety incident reports; quarterly aggregates (Title IV)
PCAOB-style audit regime	Independent audit of issuer statements	Accredited third-party audits of frontier safety claims (Title V)
Examination & enforcement	Inspections, subpoenas, civil penalties, bars	Examination of registered developers; cease-and-desist; tiered civil penalties (Title V)
Whistleblower program	Dodd-Frank awards of 10–30% of large sanctions	Identical structure for AI safety whistleblowers (Title V)
SRO supervision (FINRA)	Industry self-regulation under Commission oversight	Registered AI standards SRO; Commission approves rules and hears appeals (Title VI)
National market system	Uniform federal rules preempting conflicting state regimes	Partial preemption of state development-level mandates for compliant registrants (Title VII)

D. The honest limits of the analogy

Credibility requires acknowledging where the analogy strains, because each strain drove a design choice in the framework. **(1) Securities disclosure is disciplined by investor litigation; AI safety disclosure has no equivalent private enforcement backstop.** The framework compensates with stronger public enforcement — examination authority and a whistleblower program calibrated to surface what markets cannot — rather than by creating a new private right of action. **(2) Capability evaluation is younger than accounting.** GAAP took decades; AI evaluation science is under a decade old. The framework therefore mandates disclosure *against the developer's own stated safety framework* (an approach California pioneered) while directing

the Commission, with NIST and CAISI, to converge standards over time — and it schedules threshold recalibration biennially rather than freezing 2026 assumptions into statute. **(3) The technology moves faster than notice-and-comment.** Hence the SRO title: technical standards can iterate at industry speed inside a supervised body, with the Commission holding approval and veto authority — the division of labor that has let FINRA rules evolve continuously for decades under statutory oversight.

IV. The Bipartisan Case

A durable AI statute must give each governing coalition its principal objective while denying neither its non-negotiables. The framework is engineered around that trade, and the polling suggests the electorate has already made it: support for a new federal AI oversight agency stands at **80 percent among Democrats and 84 percent among Republicans** (Public Citizen, April 2026); 84 percent of Democrats and 83 percent of Republicans say companies should not build smarter-than-human systems until they can demonstrate control; and more than 60 percent of both parties want government rather than companies setting safety standards. Republican voters simultaneously report higher favorability toward AI as a technology — enthusiasm for the technology and demand for rules of the road are not in tension; they are the same instinct that built American capital markets.

What conservatives get

- **National uniformity with substance.** Preemption of conflicting state development-level mandates for compliant registrants — the outcome two years of moratorium attempts failed to deliver, achievable only through the exchange Congress has signaled it will accept.
- **Disclosure, not licensing.** No federal product approval, no permission slips to innovate. The default is deployment; the agency bears the burden, on the record, on a clock, under judicial review.
- **Markets over mandates.** Standardized disclosure lets insurers, enterprise customers, and investors price AI risk — the mechanism conservatives have long preferred to command-and-control.
- **Fee funding, small-government form.** A lean commission (initial staffing in the hundreds, not thousands), majority fee-funded like the SEC, with a statutory sunset forcing Congress to reauthorize or redesign within seven years.
- **Innovation in statute.** A sandbox modeled on the Cruz SANDBOX Act and Texas TRAIGA, small-entity exemptions, and capability-based treatment that never penalizes open-weight release as such.
- **National security visibility.** Statutory authority for the pre-deployment frontier access the June 2026 executive order could only request as a favor.

What progressives get

- **A real regulator, not a pledge registry.** Examination authority, subpoena power, civil penalties, stop-order authority for defined catastrophic risks, and an enforcement division — the difference between the 1933 FTC arrangement and the 1934 Commission.
- **Whistleblower protection with teeth.** Dodd-Frank-scale awards and anti-retaliation protections for AI lab employees, extending nationally what California SB 53 began.
- **Incident transparency.** Mandatory 72-hour reporting of critical safety incidents, publicly summarized — ending the era in which the public learns of model failures from lawsuits and journalists.
- **Preemption that is partial, earned, and revocable.** States keep consumer protection, civil rights, tort, child safety, and criminal law in full; preemption attaches only to registrants in good standing and lapses for developers out of compliance.
- **Civil-society standing.** Public comment on SRO rules, published enforcement outcomes, an advisory committee with public-interest seats, and annual public reporting.

Why now

Windows like this close. The 2026 election cycle has drawn super-PAC money on all sides of AI policy; each additional state statute raises the cost of the eventual federal settlement; and every month of the voluntary-access regime normalizes the idea that frontier oversight is a courtesy rather than a law. The parties' proposals are converging — the Obernolte–Trahan framework, the Hawley–Blumenthal oversight blueprint, the administration's FINRA-adjacent explorations, and industry's own supervised-standards proposals are variations on a single missing institution. The remaining task is engineering, not persuasion.

V. The Legislative Framework: Section-by-Section

What follows is the proposed architecture of the **American Artificial Intelligence Leadership and Accountability Act of 2026**, organized as bill titles with the key operative provisions of each. It is drafted as a framework for legislative counsel, not as bill text.

Sec. 2 — Findings

The Act's findings section should establish, at minimum:

- (1) Artificial intelligence systems are instrumentalities of and substantially affect interstate commerce, and the most capable systems are developed by a small number of identifiable firms whose products are deployed nationally and globally.
- (2) Forty States enacted artificial intelligence legislation in 2025, and four States have enacted materially divergent frontier-model regimes, creating compliance burdens on interstate commerce and uneven protection for the public.
- (3) The Senate, by a vote of 99–1, declined to preempt State artificial intelligence laws absent a substantive Federal framework.
- (4) No Federal agency possesses the statutory authority, technical staff, or examination capacity to verify the safety representations of frontier artificial intelligence developers.
- (5) Voluntary arrangements for pre-deployment government visibility into frontier systems, while constructive, are neither uniform nor durable.
- (6) Disclosure-based oversight of securities markets under the Securities Act of 1933 and the Securities Exchange Act of 1934 strengthened American capital markets and made them the deepest and most trusted in the world, without Federal merit approval of investment products.
- (7) A substantial majority of Americans of both political parties supports the establishment of a Federal entity to oversee the safe development of artificial intelligence.

Title I — Establishment of the Artificial Intelligence Oversight Commission

- **Composition (Sec. 101).** Five commissioners appointed by the President with Senate advice and consent; no more than three from one political party; staggered five-year terms; Chair designated by the President. Removal protections identical to the SEC's. Quorum of three.
- **Expertise requirement (Sec. 102).** At least two commissioners with substantial technical experience in machine learning or computer security; at least one with experience in a regulated industry subject to Federal disclosure requirements.
- **Offices (Sec. 103).** Statutory offices: (a) Office of Model Evaluation — the technical core, with excepted-service hiring and pay authority matching Federal financial regulators (without which the Commission cannot compete for talent and the entire design fails); (b) Office of Examinations; (c) Division of Enforcement; (d) Office of Innovation and Small Business, with a statutory advocate reporting annually to Congress; (e) Chief Economist, required to publish cost-benefit analysis for every major rule; (f) Inspector General.

- **Interagency relationships (Sec. 104).** NIST retains standards development; CAISI's national-security evaluation functions support the Commission under a mandatory interagency agreement; sector regulators (FDA, banking agencies, FAA, etc.) retain full authority over AI use within their jurisdictions, with the Commission owning model-level developer oversight — a horizontal/vertical division mirroring the SEC's coexistence with banking and insurance regulators. Classified-capability matters are coordinated with the NSA, CISA, and the Intelligence Community.
- **Funding (Sec. 105).** Registration, annual, and filing fees scaled to covered AI revenue, set to recover a target share of the Commission's budget (Congress may target 60–75 percent at maturity); authorization of appropriations of \$250 million per year for fiscal years 2027–2029; fees from small entities capped at nominal levels.

Title II — Registration and Disclosure

The tier structure is the framework's load-bearing wall: obligations scale with capability and scale, and the floor is high enough that the overwhelming majority of the AI economy is untouched.

Tier	Who is covered	Obligations
Tier 0 — Exempt	All developers below Tier 1 thresholds; academic and government research; open-source contributors; deployers and users as such	No registration, no disclosure, no fees. Express statutory exemption.
Tier 1 — Covered developers	Developers with a training run above 10^{25} FLOP and annual covered-AI revenue above \$100 million (both prongs required)	One-time registration; standardized annual disclosure report (capabilities, evaluation results, safety and security practices, aggregate compute); truthfulness obligations with penalties for material misstatement.
Tier 2 — Frontier developers	Developers with a training run above 10^{26} FLOP and annual covered-AI revenue above \$500 million (aligned with California SB 53), or designated by Commission rule on defined dangerous-capability findings, subject to judicial review	All Tier 1 duties, plus: pre-deployment Frontier Model Safety Case (Title III); 72-hour incident reporting (Title IV); periodic third-party audit and Commission examination (Title V); security and insider-risk program requirements.

- **Biennial recalibration (Sec. 203).** The Commission must revisit thresholds every two years by rule, on the record, so statute does not fossilize 2026's compute economics. Thresholds may move in either direction; algorithmic-efficiency gains are an explicit statutory factor.

- **Open-weight neutrality (Sec. 204).** Coverage turns exclusively on capability and scale, never on release modality. Open-weight publication is neither penalized nor exempted as such; the safety case for an open-weight frontier model addresses irreversibility of release, as any honest safety case must.
- **Content neutrality (Sec. 205).** Nothing in the Act authorizes the Commission to regulate the viewpoint, opinion, or lawful expressive content of model outputs. Disclosure obligations reach safety-relevant facts — capabilities, incidents, safeguards — not speech. This section is both a First Amendment firewall and a bipartisan one.

Title III — Frontier Model Safety Review

- **The filing (Sec. 301).** Before deploying a covered frontier model (or upon a defined major capability update), a Tier 2 developer files a Frontier Model Safety Case: model description; dangerous-capability evaluation results (CBRN uplift, offensive cyber, deception and autonomy, self-exfiltration); mitigations and their measured effectiveness; security posture against theft of weights; and the developer's own deployment conditions. Classified annexes are permitted; trade secrets are protected under FOIA Exemption 4 with criminal penalties for staff disclosure.
- **The clock (Sec. 302).** The filing becomes effective 90 days after submission unless the Commission, by recorded vote, issues a deficiency notice specifying what is missing. Deployment may proceed upon effectiveness. This is the securities registration mechanic: the government must act, on the record, on time — or the market proceeds.
- **Stop orders (Sec. 303).** The Commission may suspend effectiveness only upon a finding, by clear and convincing evidence, of a substantial and unmitigated risk of catastrophic harm (defined in Sec. 901 with specificity: mass casualties, critical-infrastructure incapacitation, or weapons-of-mass-destruction uplift). Stop orders receive expedited review in the D.C. Circuit within 60 days. Misuse of the power is checked by the evidentiary standard, the clock, and the court.
- **Government evaluation access (Sec. 304).** Codifies and makes uniform the pre-deployment evaluation access the June 2026 executive order requests voluntarily: secure, time-limited, read-only evaluation access for the Commission and designated national-security evaluators, with statutory protection of model weights and liability protection for developers for the access itself.

Title IV — Incident Reporting

- **Critical safety incidents (Sec. 401).** Report to the Commission within 72 hours of determination: incidents involving death or serious injury proximately involving a covered model; loss of model weights; discovery of a dangerous capability materially beyond the filed safety case; autonomous behavior defeating containment or oversight controls; and large-scale generation of prohibited content defeating safeguards. Quarterly aggregate reports cover lesser categories.
- **Safe harbor (Sec. 402).** Good-faith reports are inadmissible as an admission of liability in private litigation, and timely reporting is a mitigating factor in any Commission enforcement

— the aviation-safety insight (ASRS) grafted onto the 8-K: the system must make honesty cheaper than concealment.

- **One filing, national effect (Sec. 403).** A compliant Federal incident report satisfies any State incident-reporting requirement for the same incident; the Commission shares reports with relevant States and sector regulators through a secure clearinghouse.

Title V — Examinations, Audits, Enforcement, and Whistleblowers

- **Examinations (Sec. 501).** Periodic and for-cause examination authority over registered developers, scoped to verification of filings and incident reports — records, personnel interviews, and evaluation reproduction; not open-ended fishing expeditions.
- **Third-party audit (Sec. 502).** The Commission accredits independent AI audit firms and sets audit standards (with NIST); Tier 2 safety cases require an independent audit opinion within 18 months of the audit regime's establishment — seeding the AI equivalent of the accounting profession, a new American professional-services export.
- **Enforcement (Sec. 503).** Cease-and-desist authority; civil penalties in three tiers rising to the greater of \$10 million or 4 percent of annual covered-AI revenue for knowing material misstatements in a safety case or willful failure to report a critical incident; officer-and-director-style bars from safety-certification roles for individuals who knowingly certify false filings. All penalty actions triable de novo in Federal district court at the respondent's election, honoring *Jarkesy*.
- **No new private right of action (Sec. 504).** The Act creates no private cause of action and displaces none that exists under other law. Accountability to the public runs through the Commission; accountability to injured parties runs through the ordinary law of the States, which Title VII preserves.
- **Whistleblowers (Sec. 505).** Awards of 10–30 percent of monetary sanctions exceeding \$1 million for original information; anti-retaliation protections enforceable in district court; a protected channel that overrides contractual nondisclosure and non-disparagement clauses for safety-relevant disclosures to the Commission — nationalizing protections California began and closing the NDA gap that AI-lab employees have publicly identified.

Title VI — Self-Regulatory Organization

- **Registration (Sec. 601).** The Commission may register one or more AI standards SROs — industry-funded bodies that draft technical standards, evaluation protocols, and conduct rules for members. Every SRO rule requires Commission approval after public comment; the Commission may abrogate or amend SRO rules; SRO disciplinary actions are appealable to the Commission. This adopts the structure reportedly favored within the administration and proposed publicly by industry leadership in July 2026 — FINRA under the SEC, not FINRA instead of the SEC — and it answers the pace problem: standards iterate at industry speed inside a body the law supervises.
- **Guardrails (Sec. 602).** SRO governance must include independent public members; SRO membership may not be conditioned on anticompetitive terms; the antitrust laws apply except as to conduct required by an approved rule.

Title VII — Relationship to State Law and Other Authorities

- **Partial preemption, earned (Sec. 701).** For a registered developer in good standing, State laws that regulate the development, training, or pre-deployment safety documentation of covered models — the development layer — are preempted as to that developer. Compliance with Federal disclosure and incident-reporting satisfies equivalent State requirements (California SB 53, the New York RAISE Act, and successors).
- **What States keep (Sec. 702).** Express savings clause: nothing preempts State consumer-protection, civil-rights, tort, contract, child-safety, criminal, election, or professional-licensing law, State authority over State procurement, or State laws governing the use of AI within the State. Colorado-style deployment rules, Illinois's therapy-AI rules, and Utah's disclosure rules are untouched. State attorneys general may bring civil actions in Federal court to enforce Titles II and IV against developers operating in their States, on parity with State enforcement of Federal consumer statutes.
- **Preemption that lapses (Sec. 703).** Preemption attaches only while the developer is registered and in good standing; a developer under a final order for willful violations loses it. GAO reports to Congress on the preemption balance at years three and six. The moratorium debate ends the only way it durably can: preemption as the reward for supervised transparency, not a gift.

Title VIII — Innovation, Small Business, and Accountability

- **Regulatory sandbox (Sec. 801).** A Commission-administered sandbox granting time-limited, conditioned waivers of Commission rules for qualifying deployments, on the SANDBOX Act and TRAIGA models, with published outcomes.
- **Small-entity protection (Sec. 802).** Statutory exemption below Tier 1 thresholds (not merely regulatory grace); the Office of Innovation and Small Business must certify the small-entity impact of every major rule; expedited no-action-letter process.
- **Research exemptions (Sec. 803).** Academic, nonprofit, and Federal research are outside the registration regime; safe harbor for good-faith red-teaming and safety research on covered models, overriding contrary terms of service — the security-research lesson learned from the CFAA, applied prospectively.
- **Reauthorization sunset (Sec. 804).** The Act's regulatory titles (II–VII) sunset seven years after enactment absent reauthorization, with GAO comprehensive review at year five. An institution this novel must return to Congress on the record.
- **Reports (Sec. 805).** Annual public report on the state of frontier AI safety; annual innovation-impact report; biennial threshold rulemaking record; public enforcement statistics.

Title IX — Definitions and General Provisions

Key defined terms — *artificial intelligence model*, *covered developer*, *frontier model*, *deploy*, *critical safety incident*, *catastrophic harm*, *dangerous capability*, *covered AI revenue* — are sketched in Appendix A. Definitional drafting should track California SB 53 and the Obernolte–Trahan draft

where possible: alignment lowers transition costs for the developers already complying with those regimes and eases the preemption bargain for the States that wrote them.

VI. Anticipated Objections and Responses

“A new agency will smother a strategic industry in red tape.”

The framework touches, by design, on the order of a dozen firms at Tier 2 and a somewhat larger set at Tier 1. Everyone else — every startup, every deployer, every researcher, every open-source contributor — is statutorily exempt. The binding obligations are to tell the truth in standard formats and to report catastrophic-category incidents; the deployment default belongs to the market, with the agency on a clock. The genuinely smothering scenario is the status quo: four state frontier regimes and counting, discovery-driven disclosure through mass tort litigation, and a Brussels default. The SEC precedent is instructive — mandatory disclosure did not shrink American capital markets; it made them the world's deepest.

“This entrenches incumbents who can afford compliance.”

Regulatory capture is a real failure mode, and the framework treats it as a design constraint rather than an afterthought: dual thresholds (compute *and* revenue) keep pre-revenue startups out even as compute costs fall; the biennial recalibration is bidirectional; the Office of Innovation and Small Business has a statutory advocate; fees are progressive with revenue; the sandbox and no-action process create lawful fast paths; and the sunset forces a public re-justification within seven years. Note also what incumbents actually lobby for: uniform preemption without institutional oversight. This framework makes them earn it.

“Government cannot evaluate frontier models; the talent is all private.”

Partly true today, and the framework is built for it. Verification leans on a tripod: the developer's own evidenced filings, accredited third-party audit, and a Commission technical office with financial-regulator pay authority — supplemented by CAISI's existing evaluation capacity under mandatory interagency agreement. The SEC does not out-model Goldman Sachs; it makes Goldman sign, show its work, and answer for lies. The same asymmetry suffices here. And the whistleblower title recruits the one group that always knows the truth: the labs' own staff.

“The stop-order power will be abused to block disfavored companies or viewpoints.”

The power is confined by an evidentiary standard (clear and convincing), a defined and narrow harm category (mass casualties, critical-infrastructure incapacitation, WMD uplift), a recorded multi-member vote on a bipartisan commission, expedited Article III review on a 60-day clock, and a content-neutrality section that strips the Commission of any authority over model viewpoint or lawful expression. No existing state statute, and no executive order, offers frontier developers this much procedural protection — today a single state attorney general or a DOJ task force can move against a developer with none of these constraints.

“Why not simply let the states experiment?”

Because the experiment has run, and its results are in Section II: forty states, four incompatible frontier regimes, effective dates that moved after enactment, and a federal executive suing the laboratories of democracy for experimenting. Development-level rules for models trained once and deployed everywhere are a textbook interstate-commerce problem — which is precisely why this framework preempts *only* the development layer and leaves states every tool they have ever used to protect their citizens from harmful *uses*.

“Why not wait for the technology to settle?”

Congress in 1933 did not wait for securities to settle; it built an institution able to learn faster than fraud. Waiting has a price schedule: each Congress that passes without a framework adds state statutes to unwind, mass-tort baselines to litigate, voluntary arrangements to formalize, and foreign standards to inherit. The framework's own humility mechanisms — biennial thresholds, SRO iteration, GAO review, sunset — are how a statute stays current. Absence of a statute is not neutrality; it is a decision to let the least accountable actors set the norms.

VII. Implementation, Resources, and Timeline

The Commission should launch lean and earn scale. For reference, the SEC operates with roughly 4,500–5,000 staff and a budget near \$2.4 billion, overseeing tens of thousands of registrants. The AIOC's registrant population at launch is measured in dozens. A credible opening posture is 300–500 staff and \$150–250 million annually — comparable to a mid-sized financial regulator — weighted heavily toward the Office of Model Evaluation, reaching majority fee funding by year three. The June 2026 executive order's benchmarking process, CAISI's evaluation teams, and existing NIST standards work provide a running start; the Act should authorize personnel details and non-reimbursable transfers from those programs during standup.

Milestone	Action
Enactment + 180 days	Commissioners nominated and confirmed; Chair designates acting office heads; interagency agreements with NIST/CAISI executed; interim registration form published.
+ 12 months	Tier 1 registration opens; initial disclosure rules final; incident-reporting rule final and the 72-hour duty operative; whistleblower office open.
+ 18 months	Title III safety-case filings begin for new frontier deployments (models already deployed file retrospective safety cases within a further 12 months); preemption under Title VII attaches for registrants in good standing from first compliant filing.
+ 24 months	Audit-firm accreditation standards final; SRO applications accepted; sandbox operational.
+ 36 months	First audited safety cases; first biennial threshold rulemaking; GAO first

Milestone	Action
	preemption report; majority fee funding achieved.
Year 5 / Year 7	GAO comprehensive review (year 5); reauthorization vote before sunset of Titles II–VII (year 7).

Preemption timing is the political keystone: it attaches developer-by-developer upon compliant registration, so the benefit to industry arrives exactly as fast as the transparency arrives to the public — neither side extends credit.

VIII. Conclusion

The question before Congress is not whether artificial intelligence will be governed. It is being governed now — by four state legislatures, a DOJ litigation task force, plaintiffs' lawyers, foreign parliaments, and the self-restraint of a dozen companies racing one another. The question is whether the United States will govern it the way it governed securities: with an expert, bipartisan, independent institution that polices honesty rather than picking products, that lets markets move by default and government act on the record, and that turned American disclosure standards into the world's benchmark for ninety years.

Both parties have already rejected the alternatives — preemption without substance, and a patchwork without end. Four-fifths of both parties' voters are waiting at the destination. The Securities Exchange Act of 1934 passed a divided Congress because it gave markets what they need to grow and the public what it needs to trust them. Its successor for artificial intelligence is available on the same terms. Congress should take them.

Appendix A — Key Definitions (Drafting Sketch)

- **Artificial intelligence model.** A machine-based system trained on data that generates outputs — predictions, content, recommendations, or actions — for explicit or implicit objectives. (Track the OECD/NIST formulation used in existing Federal drafts.)
- **Covered developer (Tier 1).** A person that has trained an AI model using computing power greater than 10^{25} integer or floating-point operations and had covered AI revenue exceeding \$100,000,000 in the preceding fiscal year. Both prongs required; controlled-group aggregation rules apply.
- **Frontier developer / frontier model (Tier 2).** A covered developer with a model trained above 10^{26} operations and covered AI revenue above \$500,000,000 (harmonized with California SB 53), or a developer designated by rule upon defined dangerous-capability findings, subject to judicial review under the APA.
- **Deploy.** To make a model available for use by any person other than the developer, whether by interface, API, or release of weights; internal use above defined scale thresholds constitutes deployment for Title III purposes.
- **Critical safety incident.** An incident enumerated in Sec. 401: death or serious physical injury proximately involving a covered model's output or action; unauthorized access to or exfiltration of frontier model weights; discovery of a dangerous capability materially exceeding the filed safety case; model behavior defeating containment, shutdown, or oversight controls; or mass generation of prohibited content defeating deployed safeguards.
- **Catastrophic harm.** Death or serious injury to more than 50 persons, damage exceeding \$1,000,000,000, or incapacitation of critical infrastructure, proximately resulting from a covered model's capabilities; and any material uplift toward chemical, biological, radiological, or nuclear weapons development beyond what is achievable from public sources. (Numeric thresholds harmonized with California SB 53 and the New York RAISE Act to ease the preemption transition.)
- **Dangerous capability.** A model capability that materially enables catastrophic harm, including defined categories of CBRN uplift, autonomous offensive cyber operation, large-scale deception or manipulation of persons, and autonomous replication or exfiltration.
- **Covered AI revenue.** Revenue attributable to the development, licensing, or deployment of covered models, under Commission allocation rules; the fee base and the threshold prong.

Appendix B — Sources and Further Reading

Citations current as of July 28, 2026. Items marked with an asterisk rest on press reporting of non-public deliberations and should be characterized as “reported” if quoted.

Congressional materials

- Great American Artificial Intelligence Act of 2026, discussion draft (Reps. Obernolte and Trahan), June 4, 2026; coverage: Roll Call, “Bipartisan AI draft proposes three-year

preemption of state laws” (June 4, 2026); Tech Policy Press, “Unpacking the Great American Artificial Intelligence Act of 2026.”

- GUARD Act, S. 3062, 119th Cong. (Hawley–Blumenthal); advanced by Senate Judiciary Committee unanimously, April 30, 2026 (Covington, Global Policy Watch, May 2026).
- SANDBOX Act, S. 2750, 119th Cong. (Cruz); Senate Commerce Committee materials, September 2025.
- AI Foundation Model Transparency Act of 2026, H.R. 8094; Federal AI Risk Management Act of 2026, H.R. 8819; Future of Artificial Intelligence Innovation Act of 2026, S. 3952 (congress.gov).
- Senate vote striking the state-AI moratorium from the 2025 reconciliation act, 99–1 (July 2025); omission of the moratorium from the FY2026 NDAA (StateScoop, December 2025).
- Blumenthal–Hawley Bipartisan Framework on Artificial Intelligence Legislation (September 2023) — the earlier congressional blueprint for an independent AI oversight body with registration and audit authority.

Executive branch

- America’s AI Action Plan, “Winning the Race” (White House, July 23, 2025), and accompanying executive orders of the same date.
- Executive order on national AI policy uniformity establishing the DOJ AI Litigation Task Force (December 2025; reported as E.O. 14365 — verify number and date against the Federal Register before quoting).
- Executive order, “Promoting Advanced Artificial Intelligence Innovation and Security” (June 2, 2026): classified frontier benchmarking; voluntary 30-day pre-release access; express disclaimer of licensing authority (analysis: Skadden, June 2026).
- Center for AI Standards and Innovation (CAISI), Department of Commerce — renaming and mission refocus, June 2025.
- * Press reporting on Treasury Secretary Bessent’s proposal for a FINRA-style AI oversight body connected to the SEC (July 2026).

State law

- California SB 53, Transparency in Frontier Artificial Intelligence Act (effective January 1, 2026); AB 2013; SB 243. Analyses: White & Case (October 2025); Brookings, “What is California’s AI safety law?”
- Colorado SB 24-205 (2024), delayed by SB 25B-004 (August 2025) and rewritten by SB 26-189 (May 14, 2026; effective January 1, 2027).
- New York RAISE Act (signed December 19, 2025; chapter amendment March 27, 2026; effective January 1, 2027).
- Texas TRAIGA (signed June 22, 2025; effective January 1, 2026). Illinois PA 104-0054 (2025); Illinois HB 3773 (effective 2026); Utah AI Policy Act (2024) and 2025 amendments.
- Volume data: NCSL and Future of Privacy Forum state-AI legislation trackers (2025–2026): 40 states enacting in 2025; 1,561 bills in 45 states by March 2026.

Incidents, enforcement, and public opinion

- Character.AI and Google settlements of minor-harm suits announced January 7, 2026 (CNN; CNBC). *Raine v. OpenAI* (August 2025) and seven related California suits (November 2025).
- TAKE IT DOWN Act, Pub. L. 119-12 (signed May 19, 2025); FTC enforcement of the removal mandate from May 19, 2026.
- Grok/X image-generation incident (December 2025–January 2026) and resulting investigations (Lawfare; Privacy International; volume figures are advocacy estimates).
- Public Citizen polling memo (April 2026): 80% of Democrats and 84% of Republicans support creation of a new federal AI agency; 85/84/81% (D/I/R) support mandatory AI-content labeling.
- NBC News polling (2026): 84% of Democrats and 83% of Republicans oppose building smarter-than-human AI absent demonstrated control; majorities of both parties prefer government-set safety standards.

International

- EU AI Act implementation timeline and the Digital Omnibus amendments (European Parliament approval June 16, 2026; Annex III obligations postponed to December 2027; Article 50 transparency duties effective August 2, 2026).
- China: generative-AI measures (2023) and AI content-labeling measures (effective September 1, 2025). United Kingdom: sectoral approach; no AI statute as of mid-2026.

Scholarship and proposals

- Andrew Tutt, “An FDA for Algorithms,” 69 Admin. L. Rev. 83 (2017).
- Daniel Carpenter et al., “An FDA for AI? Pitfalls and Plausibility of Approval Regulation for Frontier Artificial Intelligence,” AIES (2024).
- Demis Hassabis, proposal for a supervised “Frontier AI Standards Body” (July 2026; Fortune, July 21, 2026).
- Securities Exchange Act of 1934, 15 U.S.C. § 78d (Commission composition: five members, staggered terms, maximum three of one party).

About the Publishers

The Align AI Foundation

The Align AI Charitable Foundation is a philanthropic funder of artificial intelligence safety and alignment research, working to advance *safe, aligned artificial intelligence for the benefit of humanity* and to keep humanity in the loop as ever more powerful systems are built. Its priorities span interpretability research that makes model decision-making legible, secure environments for evaluating systems before real-world deployment, and research into deception risk — ensuring models stay honest about their own capabilities — alongside harms already in evidence: misinformation, algorithmic bias, workforce displacement, and erosion of privacy.

Its grantmaking is theory-agnostic and evidence-driven, funding the most promising research regardless of origin or school of thought. It offers full-stack support — engineering and compute alongside capital — multi-year commitments, and transparent impact measurement. The foundation accepts no corporate revenue and holds no commercial interest in the industry it studies; its only obligation, in its own words, is to the mission of safe and beneficial AI for all of humanity. Founded by Gary Kucher.

alignai.org

The Geopolitical Insight and Education Foundation

The Geopolitical Insight and Education Foundation (GIEF) is an independent policy and research think tank focused on stability and resilience in our democracies. Founded in 2024 to examine the most pressing geopolitical issues of our time through rigorous, evidence-driven research and effective advocacy, GIEF treats democratic stability as both the central problem of our time and the enabling condition for addressing every other. Its research identifies three principal drivers of instability — technological acceleration, systemic stress, and geopolitical fragmentation — and how each affects democracies and their institutions.

GIEF works through a research–policy–advocacy loop, producing long-form reports alongside accessible policy commentary; its method is data-driven and nonpartisan — it follows where the evidence leads. Programs include *Responsible AI in Science*, on scientific integrity in academic publishing; *Drivers of Instability*; *European Resilience*; and *AI on the Ballot*, a 2026 midterm series on AI as a voter priority. GIEF is a registered 501(c)(4) nonprofit, fully independent — not funded by, or beholden to, any government, political party, or interest.

giefoundation.net

Correspondence: Geopolitical Insight and Education Foundation, 1717 N Street NW, Suite 1, Washington, DC 20036 · info@giefoundation.net · giefoundation.net · alignai.org

A working draft for congressional consideration and public comment, representing the joint views of the publishing organizations. Not legal advice. Citations current as of July 28, 2026.