

LUMENFORGE ADVISORS, LLC

Policy Analysis · Frontier AI Governance

They Built a Club and Called It Safety

How OpenAI, Anthropic, and Google's safety rules also happen to lock out everyone who isn't already inside. A fact-checked working assessment

Dan Bond · Founder & Principal AI Advisor

21 September 2026 · Version 1.3

BOTTOM LINE

The incentive is real. A secret three-firm cartel is not the best description of what we can see. OpenAI, Anthropic, and Google have converging commercial reasons to prefer high-threshold, high-fixed-cost federal rules they can already staff, plus pressure on open-weight and Chinese models that threaten closed-API economics. That is parallel self-interest with occasional public coordination, not a proven conspiracy. The risk to startups sits in the design of the instrument, not in a signed pact. Mixed motives in the same room is the honest read: some of the people writing these rules believe the models are dangerous; some would like the next lab to need a license they already hold. Both can be true at once.

What This Means For Your Business

The one-page version. No Washington jargon, no legal terms, written for anyone running a business in the Carolinas who does not have time for the other eleven pages right now.

Bottom line: the three biggest AI companies are helping write the rules that will govern AI in this country, and the rules on the table happen to be the easiest ones for the three biggest AI companies to follow. You do not have to believe that is a plan to see the effect: whatever comes out of Washington will probably assume you already have a compliance team, a lawyer on retainer, and a six-figure AI budget. Most businesses in York, Gaston, and the Charlotte region do not.

What we're seeing: Anthropic and OpenAI more than quadrupled their federal lobbying spend in a single year and both opened permanent Washington offices. The rules they are pushing ("independent testing," "licensed auditors," high-cost-threshold coverage) sound like basic safety. They are also, not by accident, the kind of rules that only a company already running a security team and a legal department can absorb without blinking. A 40-person shop cannot.

Where this matters to you:

- You almost certainly will not be regulated directly. The federal bill in play is aimed at a handful of the largest AI developers, not at the businesses using AI tools day to day.
- You will still feel it, through your vendors and your customers. If your AI tool of choice must restructure how it operates, or gets priced differently, that cost typically lands on you. And larger customers, government agencies, banks, insurers, tend to copy whatever "compliant" looks like at the top, even before any law makes them.
- Do not build your whole operation on one AI provider. If new rules squeeze one company's model or pricing, you do not want your business to go down with it. A second option, including an open-source path, is insurance, not ideology.
- Write down your own AI rules now, before someone else's version gets handed to you. A one-page policy on which tools are approved, who reviews what the AI produces, and how you handle sensitive data costs you an afternoon today. Retrofitting one under pressure, from a customer, an insurer, or a new law, costs a lot more.

Net: this fight is not about you yet. But you will get handed the bill for however it settles, without having had a seat at the table where the rules got written. The businesses that come out ahead will be the ones who already had their own house in order, not the ones scrambling to catch up.

If you want help building that one-page AI policy before it becomes someone else's requirement, that is what LumenForge's AI Governance Blueprint is for.

1. Purpose, scope, and how to read this paper

This paper examines a specific claim: that OpenAI, Anthropic, and Google are using federal AI regulation (framed as safety) as a tool to raise barriers to entry and reduce the chance that a later startup outpaces them.

Scope is deliberate. xAI / X is excluded. That firm's public posture on state patchworks, EU-style process, and default-legal environments does not sit in the same camp. Including it blurs the analysis. The three firms here share closed-model economics, large compliance staffs, and an active Washington presence.

That line gets one real stress test, in Section 2.5: xAI turns up as a co-defendant in the lawsuit this paper treats as the clearest sign that coordination has become a legal question, because Elon Musk publicly endorsed the same essay that pulled the other three labs into it. Set against that single moment: in April 2026 xAI sued to block Colorado's AI antidiscrimination law before it took effect, the Justice Department later joined on xAI's side, and the law was stayed pending the case (Colorado Sun, "Elon Musk's xAI sues over Colorado's AI law," April 10, 2026; HR Dive, "Judge stays Colorado AI bias law," 2026). That is a company fighting to keep a state AI mandate from taking hold, not lobbying to write or ratchet one up. The exclusion above holds because of that pattern, not despite the one moment where the lines blurred.

This is not a legal brief and it is not an accusation of a completed Sherman Act violation. One consumer class action filed 18 September 2026 alleges a horizontal slowdown agreement based largely on public statements. That case is an allegation. It is evidence that the coordination question is now live. It is not a finding.

Sourcing standard

Claims in this paper are tagged by confidence:

- **Tier 1: primary or official.** Lobbying disclosures as reported by Axios and the Financial Times from Senate filings; House sponsor text on the FRONTIER Act; the filed complaint caption.
- **Tier 2: major-outlet reporting.** POLITICO, Axios, AP, FT on strategy, state campaigns, and CEO statements.
- **Tier 3: interpretive.** Economic theory (Stigler), competition papers, startup commentary. Used for mechanism, not for secret-deal claims.

Where reporting describes private lobbying, that is labeled as such. Probability scores in Section 7 are judgments, not measurements. They are there to force a clean distinction between what we can document and what we infer.

2. What we can document

2.1 The influence operation is now first-order

These three are no longer occasional visitors in Washington. In Q1 2026 Anthropic spent \$1.6 million on federal lobbying and OpenAI \$1.0 million, both records at the time, against \$360,000 and \$560,000 in Q1 2025.¹ By mid-year Anthropic had spent about \$3.53 million in the first half and OpenAI about \$2.22 million. Alphabet remains

¹Lobbying Disclosure Act filings, as reported by Axios, "Anthropic outspends OpenAI in biggest-ever lobbying quarter," April 21, 2026. <https://www.axios.com/2026/04/21/anthropic-outspends-openai-biggest-lobbying-quarter>

a larger absolute spender. Anthropic opened a dedicated D.C. office in March 2026 and tripled its public-policy staff; OpenAI's first Washington office, announced in mid-2025, was operating by early 2026 (Axios, "Anthropic ramps up D.C. presence," March 11, 2026; corrects the footnote 2/3 citations below, which cover lobbying spend, not office openings).²³

The docket is broader than "please regulate us." Filings and reporting list export controls, cybersecurity, copyright, cloud and energy infrastructure, defense procurement, and model-release rules. Safety is one track inside an industrial-policy fight.

Net: treat Washington as a product surface for these firms, not a background risk.

2.2 The architecture they keep proposing is the same shape

In July 2026 Axios reported that Demis Hassabis, Sam Altman, and Dario Amodei had converged, on the record, on a diagnosis and a family of prescriptions: independent testing of frontier models, a break from pure self-attestation, and some national or international standard that can harden into market-access rules.⁴ They disagree on the referee. Amodei has argued for something closer to an FAA with blocking power. Hassabis has argued for a FINRA-style industry body under federal oversight. Altman has argued for an IAEA-style certification forum. The compliance burden lands in the same place: labs that already have security teams, lawyers, and government relations.

2.3 They are not a bloc on tactics

OpenAI's Chris Lehane has described a "reverse federalism" play: if Congress will not write a national rule, get large states to pass similar bills the industry can live with, then treat that as a de facto national standard.⁵ Anthropic has said the opposite out loud: it is not trying to use state policy as a ceiling on federal legislation, and it has pushed tougher state bills, including being the only leading lab to endorse California's first frontier-model law (SB 53, the Transparency in Frontier AI Act).⁶ OpenAI and Google have opposed some state packages Anthropic supported. That is rivalry using regulation as a weapon, not a joint venture.

2.4 The federal instrument on the table is threshold-based

The FRONTIER Act (Reps. Trahan and Obernolte, July 2026) creates a national, risk-based framework for the most advanced models: safety frameworks, transparency, independent verification organizations (IVOs), incident reporting, and emergency authority at Commerce. Coverage is tiered by the size of the developer.⁷

²³Axios, "Anthropic ramps up lobbying spending amid AI policy fights," July 21, 2026. <https://www.axios.com/2026/07/21/anthropic-ramps-up-lobbying-spending-ai-policy-fights>

³Financial Times, "AI companies spend record sums on Washington lobbying," July 27, 2026. <https://www.ft.com/content/d8a5f95e-3b6d-463a-a848-c9ef8e2394db>

⁴Axios, "Behind the Curtain: AI godfathers converge on regulations," July 16, 2026. <https://www.axios.com/2026/07/16/ai-regulations-openai-anthropic-google>

⁵POLITICO, "Inside the next phase of OpenAI's political strategy," May 20, 2026. <https://www.politico.com/news/2026/05/20/chatgpt-state-ai-fight-00928903>

⁶POLITICO, "Anthropic has a counter to OpenAI's 'reverse federalism' play," July 15, 2026. <https://www.politico.com/news/2026/07/15/inside-anthropics-state-by-state-plan-to-ratchet-up-ai-rules-00998415>

⁷Office of Rep. Lori Trahan, "Trahan, Obernolte Introduce Bipartisan FRONTIER Act," July 23, 2026. <https://trahan.house.gov/news/documentsingle.aspx?DocumentID=3823>

OpenAI has publicly backed the IVO provision.⁸ Sponsor language is explicit that the bill is aimed at the handful of companies making the largest R&D bets, not at every startup. That targeting is also how a moat gets built: the club is defined as the people who already paid the buy-in.

2.5 September 2026 made coordination a legal question

On 12 September 2026 Amodei published an essay arguing the frontier should be paced, third-party evaluators should be embedded inside labs, and Washington should issue a narrow antitrust waiver so rivals can talk safety without Sherman Act risk. Altman and Hassabis publicly agreed with pieces of that within hours — and so, within hours, did two figures outside this paper’s three-firm scope: Elon Musk (“Dario is right”) and Microsoft’s Satya Nadella. A consumer class action filed 18 September in the Northern District of California, *Buist v. Anthropic PBC*, treats those public statements as a horizontal agreement to slow competing products. Note the scope tension: Buist names four defendants, not three — Anthropic, OpenAI, Google, and xAI (captioned “SpaceXAI”) — because Musk’s public endorsement pulled xAI into the same alleged agreement that Section 1 excludes from analysis. Section 1 has the other half of that picture: the same month this suit puts xAI in the room with the other three, xAI was in federal court suing to keep a state AI law from taking effect. The single legal filing anchoring this section is not confined to the three firms this paper otherwise tracks; the firm-by-firm pattern this paper tracks still is.⁹¹⁰ Amodei flagged the antitrust problem in the essay itself. Early Washington reaction to the waiver request was cold. This paper treats the filing as a fact and the conspiracy as unproven.

3. The economic logic: why the incentive exists without a conspiracy

George Stigler’s 1971 point still does the work: firms in high-fixed-cost industries often prefer regulation that raises the cost of entry, especially when the rule can be written around current practice.¹¹ Frontier training already costs hundreds of millions to billions. Adding another fixed cost (licensed auditors, safety and security protocols, embedded evaluators, incident regimes) is cheap at that scale and expensive for a 40-person shop.

Foundation-model economics already lean toward concentration: high training costs, near-zero marginal cost of inference, first-mover advantages in distribution.¹² Regulation is not required to produce a concentrated market. It can freeze one that is already forming.

Three structural facts shared by these labs make the same instrument attractive:

- The product is an API or tightly controlled weights, not a freely forked model.
- The P&L threats that move are (a) a later lab that reaches “good enough” cheaper and (b) open-weight models that collapse pricing.

⁸POLITICO, “OpenAI backs bipartisan House plan for third-party safety assessments,” September 15, 2026.

<https://www.politico.com/news/2026/09/15/openai-backs-bipartisan-house-plan-for-third-party-safety-assessments-01076588>

⁹AP News / U.S. News, “Lawsuit Says Anthropic, OpenAI, SpaceXAI and Google Made Illegal Agreement on AI Slowdown,” September 19, 2026. <https://www.usnews.com/news/business/articles/2026-09-19/lawsuit-says-anthropic-openai-spacexai-and-google-made-illegal-agreement-on-ai-slowdown>

¹⁰Complaint, *Buist v. Anthropic PBC*, No. 3:26-cv-10693 (N.D. Cal. filed Sept. 18, 2026).

¹¹George J. Stigler, “The Theory of Economic Regulation,” *Bell Journal of Economics and Management Science* 2, no. 1 (1971): 3-21.

¹²Brookings, “Market concentration implications of foundation models.” <https://www.brookings.edu/articles/market-concentration-implications->

- They already employ most of the people qualified to evaluate frontier systems. An “independent” evaluator market drawn from that pool is independent on paper.

You do not need a smoke-filled room for three firms with the same cost curve to prefer the same rule.

4. Firm-by-firm map

	Default posture	Where they fight the other two
Anthropic	Most aggressive safety-first lab. State-by-state upward ratchet. Only major lab to endorse California’s first frontier law. Large PAC spend via Public First Action.	Will accept more friction if it binds the field. Wants blocking power and a waiver for safety talks.
OpenAI	Reverse federalism: make large-state bills look alike. Waters some bills down, endorses others after amendment. Backs FRONTIER IVOs. Wants consistency and preemption more than maximum stringency.	Competes with Anthropic over who writes the template. Has historically argued some national-security issues belong in Congress, not Sacramento.
Google	Pragmatic two-track federal regime: heavier on frontier, lighter on everyday AI. Often votes with OpenAI against specific state bills. Opposed original SB 1047.	Least evangelical. Most focused on avoiding a 50-state mess and protecting the broader stack.

Historical marker, not current law: on California SB 1047 in 2024, Anthropic moved to qualified support after amendments; OpenAI and Google raised objections or opposed the original construct.¹³¹⁴ The split has persisted into 2026 state fights. Alignment is on architecture, not on who holds the pen.

5. Four channels that threaten startups

Channel 1: Thresholds as a membership card

A rule that attaches only above a revenue or compute spend line sounds targeted. In practice it freezes the current roster. Stay under the line and you are defined as not-frontier. Cross it and the first serious check goes to lawyers and auditors instead of GPUs. Sponsor claims that this “protects startups” are half right: it protects them from the statute, and it also keeps them off the field the statute governs.

Channel 2: Licensed IVOs as a de facto release gate

Commerce-licensed independent verification organizations, with a duty to flag imminent catastrophic risk and a path to emergency restriction, is a permissioning regime that does not use the word license for the model. It licenses the auditor. OpenAI’s support for this provision is rational if you already run an eval stack and already know the evaluator market. It is a different problem for a lab that does not.

¹³Carnegie Endowment, “All Eyes on Sacramento: SB 1047 and the AI Safety Debate,” September 11, 2024. <https://carnegieendowment.org/posts/2024/09/california-sb1047-ai-safety-regulation>

¹⁴Wikipedia compilation of contemporaneous industry positions on SB 1047 (cross-checked against Carnegie and contemporaneous press). https://en.wikipedia.org/wiki/Safe_and_Secure_Innovation_for_Frontier_Artificial_Intelligence_Models_Act

Channel 3: Open-weight and Chinese-model containment

This is the sharpest commercial edge. Reporting in July 2026 described OpenAI and Anthropic privately pressing Washington on Chinese open-weight models. Anthropic stayed off a multi-firm industry letter defending open weights; OpenAI signed only after the omission became public; Google signed.¹⁵¹⁶ A national-security frame around distillation is also a pricing frame around Meta, Mistral, and every fine-tune shop. Closed APIs survive that fight. Open stacks do not.

Channel 4: Uniform refusal surfaces

If every model must refuse the same class of prompts, users migrate to the best-known closed model. Startups usually live in niches large labs abandon for brand and liability reasons. Flatten the refusal surface and those niches shrink even if the safety rationale is sincere.¹⁷

A quieter fifth channel is political spend as candidate filter. Anthropic committed \$20 million to the Public First Action super PAC in February 2026 and doubled that to \$40 million for the 2026 midterms; OpenAI-linked spending in state races is being used the same way: to elevate or punish authors of particular bills.¹⁸ That is each lab buying the legislature it wants. It still raises the cost of entry for anyone not in the room.

6. What the capture thesis gets wrong

Safety concerns are not automatically a cover story. Cyber, bio, and loss-of-control questions at the current frontier are not 2019 SaaS problems. A use-only regime (fraud, employment discrimination, CSAM) leaves training-time risk on the table. Some of the people inside these buildings believe that risk is real. Discounting that is as sloppy as treating every safety sentence as a moat.

They also fight each other. Anthropic versus OpenAI on state strategy. OpenAI and Google versus Anthropic on particular state packages. Different preferred agencies. If this were a joint venture, that opposition would be theater. It does not look like theater.

Concentration is already being produced by compute, data, and distribution. Efficiency shocks that drop the resource threshold (the DeepSeek-class pattern) cut the other way. A licensing regime keyed to 2024 training costs will either miss the next wave or get rewritten by the same three labs.

Meta, Nvidia, the open-weight coalition, Y Combinator, a16z, and several model startups are not in this group. Cohere's Aidan Gomez has used the word "cartel" in public commentary about leading labs asking to set terms for everyone else. That opposition is data. If this were "the industry," it would not exist.

¹⁵Implicator / reporting on private lobbying regarding Chinese open-weight models, July 25, 2026. <https://www.implicator.ai/openai-anthropic-lobby-washington-open-weight-ai/>

¹⁶AI Weekly, "OpenAI, Anthropic Skip 25-Firm Open-Weights AI Coalition Letter," July 25, 2026. <https://aiweekly.co/alerts/openai-anthropic-skip-25-firm-open-weights-ai-coalition-letter>

¹⁷Forbes / LSE, "How AI Safety Rules Could Backfire On Competition," January 8, 2026.

<https://www.forbes.com/sites/londonchoolofeconomics/2026/01/08/how-ai-safety-rules-could-backfire-on-competition/>

¹⁸NY Post, "Anthropic spends millions in fight over AI regulations," September 19, 2026. <https://nypost.com/2026/09/19/us-news/anthropic-spends-millions-in-fight-over-ai-regulations-in-congress-as-lawmakers-allege-tech-giant-is-freezing-out-rivals/>

One more limit: OpenAI has a documented history of lobbying to narrow rules that would have treated general-purpose models as high-risk under the EU AI Act.¹⁹ “Asking for regulation” and “asking for the regulation we can operationalize” are not the same sentence. The second is the one that shows up in the record.

7. Probability matrix: likely intent

These scores are judgments against the public record as of 21 September 2026. They are not forecasts of legislative outcomes. Read “intent” as the mix of motives that best explains observed behavior, not as mind-reading.

Hypothesis	Probability	Why this score
Sincere catastrophic-risk concern is a real motive inside all three firms	High (70 to 85%)	Consistent public writing, employee letters, and willingness to accept some product friction. Does not prove it is the only motive.
Rule design is being optimized for compliance systems these labs already run	High (75 to 90%)	IVO model, SSP-style frameworks, spend/compute thresholds, embedded evaluators. This is the load-bearing claim.
Regulation is being used as a competitive weapon among the three, not only against startups	High (70 to 85%)	Opposite state strategies; split endorsements; PAC-on-PAC spending in state races.
Open-weight / Chinese-model restrictions are pursued in part to defend closed-API pricing	Medium-high (60 to 75%)	Private lobbying reports plus public absence from the open-weights letter. National-security rationale can be genuine and still convenient.
Primary intent is to freeze the market so no startup can reach the frontier	Medium-low (25 to 40%)	Too clean. They still race each other. Capture is a byproduct they will accept, not a single master plan.
A completed, illegal horizontal agreement to slow products already exists	Low-medium (15 to 35%)	Public statements and a filed complaint. No waiver granted. No adjudicated facts. Treat as live legal risk, not established fact.
An antitrust waiver for “safety conversations” will pass	Low (10 to 25%)	The ask is documented. Bipartisan allergy to competitor safe harbors is also documented.
A FRONTIER-style IVO regime becomes the federal template within 24 months	Medium (40 to 55%)	Bipartisan sponsors, industry support for the auditor piece, midterm calendar. Preemption fight is the swing variable.

NOTE: Rows 1-4 are grounded and supported by facts. Rows 5 and 6 are the ones that generate headlines and lawsuits. Do not sell them as proven.

¹⁹TIME, “Exclusive: OpenAI Lobbied the E.U. to Water Down AI Regulation.” <https://time.com/6288245/openai-eu-lobbying-ai-act>

A simple decision rule

Ask one question of any proposal: does the instrument bind conduct that causes harm, or does it bind the right to exist at the frontier?

Liability, transparency, incident reporting, and use-based rules sit on the first side. Pre-release permission, licensed auditors with blocking power, spend/compute membership cards, and open-weight containment sit on the second. The first can be safety policy. The second is industrial policy whether anyone says the words.

8. Guidance: what to watch

Do not watch rhetoric. Watch the instrument and who can operationalize it.

1. **Text of coverage thresholds.** If FRONTIER or a successor defines “very large frontier developer” by trailing spend or compute, write down who is inside the line on day one. If the line moves only upward, that is a ratchet.
2. **IVO licensing criteria and staffing.** Who gets licensed. Who funds them. Whether an adverse finding can delay release. Whether startups below the threshold still need an IVO to sell to government or regulated industries.
3. **Preemption language.** A federal floor with no ceiling leaves Anthropic’s state ratchet intact. A federal ceiling written around current lab practice is OpenAI’s reverse-federalism end state.
4. **Open-weight treatment in national-security vehicles.** Watch procurement guidance, “know your customer” on fine-tunes, and soft-law warnings that make open weights legally radioactive without a statute that says “ban.”
5. **The waiver.** Any bill that immunizes competitor-to-competitor pacing talks is the closest statutory cousin to classic collusion. Low odds. High signal if it moves.
6. **Who is missing from the coalition.** Meta, Nvidia, Microsoft-on-open-weights, and the startup/VC bloc. If they flip, the capture story gets stronger. If they stay opposed, the “industry plot” story stays weak.
7. **Efficiency shocks.** A model that reaches dangerous capability below the statutory compute line breaks the whole theory of the bill. Watch whether Congress then lowers the line (which would start taxing mid-tier labs) or ignores the miss.
8. **Discovery in Buist.** Public essays are not a conspiracy. Emails about jointly delaying a release would be. Until that record exists, do not treat the complaint as fact.

9. Implications for operators

If you run an SMB, an agency, or a mid-market builder, this fight is not abstract.

- Do not build a stack that can only be served by one licensed frontier API. Dual-source closed plus one open-weight path is risk management, not ideology.
- Assume that “frontier-only” compliance will leak into procurement even if the statute says it does not. Government and regulated buyers copy the template.

- Treat safety process as a product requirement either way. The question is whether you own the process or rent it from a lab's auditor.
- If you advise policymakers: prefer outcome liability and incident reporting over pre-clearance. Liability scales with harm. Audits scale with headcount.

10. Conclusions

Strip the romance out of both sides.

The safety camp is not a cartoon. Frontier systems raise questions that product-liability law and Section 230-era thinking do not answer cleanly. Asking for third-party testing is not, by itself, a crime against startups.

The capture camp is not a cartoon either. High-fixed-cost rules written by the only firms that can staff them, aimed at the layer where new entrants would have to compete, with a side campaign against the open-weight models that collapse pricing: that is a moat. Calling it safety does not change the industrial-policy effect.

The reasonable conclusion is mixed motives plus path dependence. These three firms will keep racing each other. They will also keep asking Washington for a rulebook that looks like the one they already use. Startups will not be the named target. They will be the people who cannot buy the compliance stack.

Watch the instrument. If the next federal bill is liability, transparency, and incident reporting, the market can live with it. If it is pre-release permission, licensed auditors with blocking power, spend thresholds, and open-weight containment, you have built a club and called it a safeguard.

That is the line. Everything else is noise.

Appendix A: Source list

Primary and major-outlet sources used in this paper. Interpretive literature is listed last.

1. Lobbying Disclosure Act filings, as reported by Axios, “Anthropic outspends OpenAI in biggest-ever lobbying quarter,” April 21, 2026. <https://www.axios.com/2026/04/21/anthropic-outspends-openai-biggest-lobbying-quarter>
2. Axios, “Anthropic ramps up lobbying spending amid AI policy fights,” July 21, 2026. <https://www.axios.com/2026/07/21/anthropic-ramps-up-lobbying-spending-ai-policy-fights>
3. Financial Times, “AI companies spend record sums on Washington lobbying,” July 27, 2026. <https://www.ft.com/content/d8a5f95e-3b6d-463a-a848-c9ef8e2394db>
4. Axios, “Behind the Curtain: AI godfathers converge on regulations,” July 16, 2026. <https://www.axios.com/2026/07/16/ai-regulations-openai-anthropic-google>
5. POLITICO, “Anthropic has a counter to OpenAI’s ‘reverse federalism’ play,” July 15, 2026. <https://www.politico.com/news/2026/07/15/inside-anthropics-state-by-state-plan-to-ratchet-up-ai-rules-00998415>
6. POLITICO, “Inside the next phase of OpenAI’s political strategy,” May 20, 2026. <https://www.politico.com/news/2026/05/20/chatgpt-state-ai-fight-00928903>
7. Office of Rep. Lori Trahan, “Trahan, Obernolte Introduce Bipartisan FRONTIER Act,” July 23, 2026. <https://trahan.house.gov/news/documentsingle.aspx?DocumentID=3823>
8. H.R. 9925, FRONTIER Act (introduced text), 119th Congress. <https://www.govtrack.us/congress/bills/119/hr9925/text/ih>
9. POLITICO, “OpenAI backs bipartisan House plan for third-party safety assessments,” September 15, 2026. <https://www.politico.com/news/2026/09/15/openai-backs-bipartisan-house-plan-for-third-party-safety-assessments-01076588>
10. AP News / U.S. News, “Lawsuit Says Anthropic, OpenAI, SpaceXAI and Google Made Illegal Agreement on AI Slowdown,” September 19, 2026. <https://www.usnews.com/news/business/articles/2026-09-19/lawsuit-says-anthropic-openai-spacexai-and-google-made-illegal-agreement-on-ai-slowdown>
11. Complaint, Buist v. Anthropic PBC, No. 3:26-cv-10693 (N.D. Cal. filed Sept. 18, 2026).
12. Carnegie Endowment, “All Eyes on Sacramento: SB 1047 and the AI Safety Debate,” September 11, 2024. <https://carnegieendowment.org/posts/2024/09/california-sb1047-ai-safety-regulation>
13. Wikipedia compilation of contemporaneous industry positions on SB 1047 (cross-checked against Carnegie and contemporaneous press). https://en.wikipedia.org/wiki/Safe_and_Secure_Innovation_for_Frontier_Artificial_Intelligence_Models_Act
14. Implicator / reporting on private lobbying regarding Chinese open-weight models, July 25, 2026. <https://www.implicator.ai/openai-anthropic-lobby-washington-open-weight-ai/>
15. AI Weekly, “OpenAI, Anthropic Skip 25-Firm Open-Weights AI Coalition Letter,” July 25, 2026. <https://aiweekly.co/alerts/openai-anthropic-skip-25-firm-open-weights-ai-coalition-letter>
16. Forbes / LSE, “How AI Safety Rules Could Backfire On Competition,” January 8, 2026. <https://www.forbes.com/sites/londonschoolofeconomics/2026/01/08/how-ai-safety-rules-could-backfire-on-competition/>
17. Brookings, “Market concentration implications of foundation models.” <https://www.brookings.edu/articles/market-concentration-implications->
18. George J. Stigler, “The Theory of Economic Regulation,” Bell Journal of Economics and Management Science 2, no. 1 (1971): 3-21.
19. AI & SOCIETY, “AI safety and regulatory capture,” 2025. <https://link.springer.com/article/10.1007/s00146-025-02534-0>
20. Bloomberg reporting via SBS, “Startups Criticize AI Safety Standards Alliance as Regulatory Barrier to Competition,” September 15, 2026.
21. POLITICO, “Top AI companies are asking for supervision,” September 15, 2026. <https://www.politico.com/news/2026/09/15/a-starting-gun-ai-leaders-push-for-independent-safety-audits-01078262>
22. TIME, “Exclusive: OpenAI Lobbied the E.U. to Water Down AI Regulation.” <https://time.com/6288245/openai-eu-lobbying-ai-act>
23. Fast Company, “Where OpenAI, Anthropic, Google, Meta, and other AI giants stand on regulation,” July 23, 2026. <https://www.fastcompany.com/91578113/where-openai-anthropic-google-meta-and-other-ai-giants-stand-on-regulation>
24. NY Post, “Anthropic spends millions in fight over AI regulations,” September 19, 2026. <https://nypost.com/2026/09/19/us-news-anthropic-spends-millions-in-fight-over-ai-regulations-in-congress-as-lawmakers-allege-tech-giant-is-freezing-out-rivals/>

25. New York Times, “California A.I. Bill Causes Alarm in Silicon Valley,” August 14, 2024.

Appendix B: What this paper does not claim

- It does not claim a proven criminal or civil conspiracy. Buist is a filed allegation.
- It does not treat xAI, Meta, Microsoft, or Nvidia as members of this coalition.
- It does not assert that catastrophic-risk arguments are false.
- Probability bands are the author’s judgment against the public record as of 21 September 2026 and should be revised as filings, bill text, and court records move.

Appendix C: Fact-check and editorial notes (21 September 2026 pass)

Every citation in this draft was checked against the underlying reporting. No source was found to be fabricated or mistitled. The corrections below are precision and scope fixes, not retractions.

One development since this draft that isn't yet reflected: OpenAI's negotiation over California's “Adam's Law” chatbot bill (reported by Fortune, September 14, 2026) is a live, concrete instance of the “waters some bills down, endorses others after amendment” pattern the firm-by-firm table already attributes to OpenAI in §4. Worth a line if you revise again before this goes out the door.

What this fact-check did not do: verify every reporter's underlying access to lobbying disclosures or court filings first-hand, or independently confirm quotes attributed inside paywalled articles beyond what secondary aggregation and outlet excerpts show. Tier 1 items (the FRONTIER Act text, the Buist docket) were checked against primary or near-primary sources; Tier 2 and 3 items rely on the outlets' own reporting, consistent with this paper's own sourcing standard in §1.

© 2026 LumenForge Advisors, LLC. Working draft for discussion. Not legal advice.