

The Architecture of Capture: How Frontier AI Models Leverage Safety and Regulation as Competitive Moats in 2026

The artificial intelligence sector in 2026 is defined by a profound paradox: the corporate entities warning most vociferously about the existential dangers of artificial intelligence are the same entities furiously racing to build it. Across the industry, leading foundational model developers—most notably Anthropic—have adopted a highly publicized public posture centered on caution, safety, and the urgent need for comprehensive government regulation. To the uninformed observer, these actions appear as an act of unprecedented corporate altruism, wherein private enterprises willingly subjugate their own profit motives and development velocity to safeguard humanity against the risks of artificial general intelligence (AGI). However, a structural analysis of the industry's policy interventions, product deployment strategies, and legislative lobbying over the past three years reveals a sharply different reality.

Beneath the rhetoric of existential risk lies a classic economic mechanism: regulatory capture. This phenomenon, famously articulated by venture capitalist Bill Gurley in 2023, occurs when an industry utilizes the coercive power of the state to protect its own incumbent interests, often under the guise of public welfare. By framing artificial intelligence as an inherently hazardous technology that requires stringent, resource-intensive oversight, incumbent frontier model developers are effectively advocating for a regulatory environment that institutionalizes their market dominance. The push for AI safety is not merely a technical or ethical framework; it is a sophisticated corporate strategy designed to erect insurmountable barriers to entry, thereby neutralizing the threat of open-source challengers and well-funded startups.

This exhaustively detailed report analyzes the 2026 artificial intelligence landscape through the lens of regulatory capture. By examining federal and state legislative battles, the evolution of proprietary safety frameworks, the restriction of advanced cybersecurity models, and the systematic suppression of open-source ecosystems, this analysis demonstrates how "safety" has been successfully weaponized as the ultimate competitive moat in Silicon Valley.

The Economic Theory of Regulatory Capture in the Intelligence Age

The foundational economic theory required to deconstruct the current artificial intelligence landscape is rooted in the concept of regulatory capture, a dynamic wherein an industry leverages state apparatuses to insulate itself from free-market competition. In the context of Silicon Valley, this phenomenon has been historically observed in sectors ranging from telecommunications to financial technology. However, in 2026, the artificial intelligence industry has perfected this model by leveraging what economists refer to as the "Bootleggers and Baptists" dynamic. In this theoretical model, regulations are successfully passed through the

combined, often unspoken, lobbying efforts of two disparate groups: the "Baptists," who desire regulation for genuine moral or safety reasons, and the "Bootleggers," who profit immensely from the market restrictions and reduced competition those regulations inevitably create.

In the contemporary artificial intelligence ecosystem, the frontier model companies play both of these roles simultaneously. They employ massive, public-facing research teams that act as the "Baptists," authoring voluminous policy papers on catastrophic biosecurity risks, autonomous alignment failures, and societal collapse to generate a persistent state of moral panic.¹ These research teams publish findings that consistently highlight the terrifying potential of the models they are currently building, ensuring that the media and regulatory bodies remain hyper-focused on catastrophic threat vectors.² Simultaneously, the executive and legal arms of these same companies act as the "Bootleggers," utilizing the resulting public anxiety to advocate for strict computing thresholds, mandatory third-party safety audits, and complex liability frameworks that are fundamentally incompatible with startup economics.⁵

The cost of compliance in an over-regulated artificial intelligence market functions as a regressive tax on innovation. When regulatory regimes mandate defense-in-depth deployment architectures, continuous red-teaming, multi-party model weight authorization, and vast legal departments to navigate conflicting state and federal laws, the absolute capital cost of bringing a foundation model to market skyrockets.⁸ An incumbent valued at \$85 billion can comfortably absorb tens of millions of dollars in compliance expenditures, treating it as a standard cost of doing business.⁵ A disruptive startup cannot. Therefore, the incumbent's calls for "pausing" development or establishing rigorous regulatory agencies serve to pull up the ladder behind them, effectively freezing the market structure at the exact historical moment they have achieved dominance.¹¹

Critics of this approach, including prominent venture capitalist Marc Andreessen, have repeatedly pointed out that the industry is experiencing a "full-blown moral panic" characterized by hysterical and irrational fear.³ Andreessen argues that artificial intelligence is simply a computer program owned and controlled by people, and that the "doomer" narrative is being actively cultivated to justify the creation of a cartel.³ This dynamic represents the realization of the exact regulatory capture scenario Bill Gurley warned about in 2023. The incumbents are not trying to stop artificial intelligence; they are trying to stop anyone else from building it.

The Legislative Battlefield: State Patchworks and the Compliance Tax

The implementation of this regulatory capture strategy relies heavily on the manipulation of the legislative process. By early 2026, the regulatory environment in the United States had fractured into an unmanageable patchwork of state-level laws, creating a chaotic compliance environment that simultaneously crushed startups and set the perfect stage for a strategic federal intervention by the incumbents.

Throughout 2025 and early 2026, state legislatures, responding to the persistent drumbeat of AI safety concerns, enacted a dizzying array of governance statutes. Colorado passed the most comprehensive state-level artificial intelligence governance law in the country, Texas enacted the Responsible AI Governance Act (TRAIGA) on January 1, 2026, and states like Illinois and New York introduced strict algorithmic accountability and biometric privacy mandates.⁶ California, maintaining its historical status as the most complex compliance environment in the nation, pushed the boundaries of legislative oversight further than any other jurisdiction.¹⁴

The evolution of California's regulatory approach perfectly illustrates the tension between market realities and regulatory capture. In 2024, California Governor Gavin Newsom vetoed the highly controversial Senate Bill 1047 (the Safe and Secure Innovation for Frontier Artificial Intelligence Models Act).¹⁵ SB 1047 would have created a draconian Board of Frontier Models, mandated independent third-party compliance audits, and required developers to install physical "kill switches" to remotely disable models.⁵ The California Chamber of Commerce and the startup ecosystem aggressively opposed the bill, arguing that holding original developers liable for third-party misuse and implementing "Know Your Customer" (KYC) requirements would devastate the state's technology sector and drive innovation overseas.⁵

However, the defeat of SB 1047 merely paved the way for a more targeted mechanism of incumbent protection: Senate Bill 53, the Transparency in Frontier Artificial Intelligence Act, which took effect on January 1, 2026.¹⁶ While characterized as a "pared-down" version of SB 1047, SB 53 represents a quintessential example of legislation that structurally favors incumbents through the explicit mechanism of computational thresholds.¹⁶

The law legally defines a "frontier model" as a foundation model that was trained using a quantity of computing power greater than 10^{26} integer or floating-point operations.¹⁷ By relying on raw computational expenditure as a proxy for catastrophic risk, the legislation establishes a distinct, mathematical demarcation line between the dominant players and the rest of the ecosystem.¹⁷ Under SB 53, developers who cross this threshold—and specifically "large frontier developers" whose corporate affiliates possess annual gross revenues exceeding \$500 million—are mandated to implement and publish a comprehensive "Frontier AI Framework".⁷ This framework requires detailed documentation of how catastrophic risks are identified, assessed, and mitigated throughout the lifecycle of the model, including cybersecurity practices and alignment with industry standards.¹⁸ Furthermore, developers must notify the California Office of Emergency Services (Cal OES) of critical safety incidents within 15 days, or 24 hours if the incident poses imminent danger.⁷ The law also enforces strict whistleblower protections, barring employers from retaliating against employees who report catastrophic risks to the Attorney General or federal authorities.⁷

For startups attempting to break into the frontier model space, these state laws represent a devastating and insurmountable headwind. Operating an AI system that is perfectly legal in Texas may suddenly trigger mandatory algorithmic audits in New York and biometric privacy violations in Illinois.⁶ The legal overhead required to navigate this environment is staggering. As

compliance costs multiply across jurisdictions, smaller entities are forced to either restrict their geographic deployment, intentionally throttle their model capabilities to stay below the 10^{26} FLOPs threshold, or abandon ambitious model training entirely.⁶

State AI Legislation (2025-2026)	Key Mechanisms & Statutory Requirements	Strategic Market Impact
California SB 53 (TFAIA)	Targets models trained on > 10^{26} FLOPs. Mandates Frontier AI Frameworks, incident reporting to OES, and strict whistleblower protections. ⁷	Cements compute expenditure as a proxy for risk. Establishes heavy documentation and legal overhead tailored exclusively for incumbents.
Colorado AI Act	Comprehensive algorithmic accountability, bias auditing, and consumer protection governance frameworks. ⁶	Elevates compliance costs significantly; delayed implementation highlights the extreme regulatory complexity facing developers. ²¹
Texas TRAIGA	Responsible AI Governance Act, enforcing state-specific transparency standards for local operations and government contracts. ⁶	Contributes to the national "patchwork" effect, necessitating bespoke regional compliance strategies that startups cannot afford. ²²
California SB 942	The AI Transparency Act requires AI providers to disclose when content is AI-generated, including through mandatory digital watermarking. ¹⁴	Adds a technical engineering burden to deployment, forcing companies to maintain complex provenance tracking systems.

The Preemption Gambit: The 2026 White House National Policy Framework

The resulting compliance nightmare at the state level provided the perfect pretext for incumbent intervention at the federal level. Having successfully established a chaotic,

fragmented regulatory landscape that startups could not navigate, the frontier model developers shifted their lobbying focus to Washington, D.C., seeking a unified federal solution that they could directly control. On March 20, 2026, the White House released its highly anticipated National Policy Framework for Artificial Intelligence.²³ Supported heavily by the major technology oligopolies, the Framework outlined legislative recommendations for Congress to establish a "unified federal approach" to AI regulation.²³

The absolute centerpiece of the White House Framework is the concept of targeted federal preemption. The administration explicitly urged Congress to adopt legislation that would broadly preempt state artificial intelligence laws deemed to impose "undue burdens" on innovation.²³ The framework's core argument is that a "patchwork of 50 different regulatory regimes" paralyzes American innovation, damages economic security, and compromises national supremacy in the ongoing technological race against geopolitical adversaries.²² To rectify this, the Framework proposes a "minimally burdensome national standard," limiting state enforcement primarily to existing, generally applicable laws protecting children, regulating internal state government procurement, and managing local infrastructure zoning for data centers.¹⁴

From the perspective of regulatory capture, the White House Framework is a strategic masterstroke by the incumbent developers. Technology giants lobbied extensively for this exact outcome because it achieves two critical, synergistic goals simultaneously.²⁸ First, it cleanly wipes the slate of stringent, aggressive state laws (such as California's transparency mandates or Colorado's bias audits) that impose genuine operational friction and legal liability on their massive deployments.²² Second, the proposed "minimally burdensome national standard" is inherently shaped by the lobbying efforts and technical input of the largest corporations, effectively allowing them to write the rules that govern their own industry.²²

Furthermore, the Framework explicitly recommends against the creation of any new standalone federal rulemaking body to regulate artificial intelligence, suggesting instead that the government support the development of sector-specific applications through existing regulatory bodies and "regulatory sandboxes".²⁵ By favoring a "lighter-touch federal approach," the framework ensures that the government lacks the centralized technical expertise required to meaningfully challenge the assertions of the frontier labs.²⁷

Policy critics and advocacy organizations, including Public Knowledge, have aggressively pointed out the fundamental flaws and industry bias within this approach. They argue that the Framework effectively asks Congress to shut down state-level democratic oversight in exchange for a federal regime that completely defers to "industry-led standards".²⁹ Because industry standards bodies—such as the Frontier Model Forum—are composed exclusively of the companies being regulated, their technical inputs are neither enforceable nor democratically accountable.²⁹ By relying on these entities, the federal government delegates the creation of the national rulebook to the monopolists.²⁹ Consequently, the industry leaders are able to craft a federal regulatory ceiling that is high enough to legitimize their own

profitable operations, but structured with enough specific infrastructural and safety prerequisites to effectively lock out open-source projects and lean startups.²⁸ The Framework's language regarding "generally applicable laws" is also predicted to spawn years of litigation, as any state consumer protection action will inevitably be met with the argument that the law is preempted because it is being applied in an AI context.²⁹

Privatizing Regulation: The Evolution of Responsible Scaling Policies

If federal and state legislation provide the external legal moat for artificial intelligence incumbents, proprietary corporate risk management frameworks provide the internal, pseudo-scientific justification for that moat. Anthropic has pioneered the utilization of the "Responsible Scaling Policy" (RSP) as a mechanism to set de facto industry standards. A detailed analysis of Anthropic's RSP, particularly its evolution from its inception to Version 3.1 in early 2026, demonstrates exactly how internal safety frameworks serve to consolidate market power and shift the burden of compliance onto competitors.

Anthropic's RSP operates on a system of tiered "AI Safety Levels" (ASL), which are explicitly modeled after the United States government's rigorous biosafety level (BSL) standards used in infectious disease laboratories.⁸ The core premise of the RSP is that as an artificial intelligence model is scaled up and crosses specific, predefined "Capability Thresholds," it automatically triggers exponentially stricter safety, security, and deployment mandates.³¹

In Version 3.1 of the RSP (effective April 2, 2026), these Capability Thresholds are rigorously defined and highly complex. For example, the "AI R&D Capability Threshold" is triggered when an artificial intelligence system demonstrates the ability to compress two years of aggregate AI progress (specifically benchmarking 2018–2024 progress) into a single year.³² Other thresholds include the "CBRN Development Threshold," which flags capabilities that could substantially uplift the development programs of moderately resourced nation-states regarding chemical, biological, radiological, and nuclear technologies, as well as the "Sabotage Risk Threshold," which monitors models for autonomous, misaligned goals.³²

When a model crosses these thresholds, it enters ASL-3 territory, which mandates a comprehensive, highly expensive "defense-in-depth" technical architecture.⁹ To legally deploy an ASL-3 model under this framework, a developer must design and implement:

1. **Tiered Access Controls:** Enhanced due diligence and compartmentalized access restrictions for trusted users, requiring deep corporate intelligence gathering on API consumers.⁹
2. **Real-Time Classifiers:** Machine learning models that analyze user inputs and generated outputs with streaming implementations designed to minimize latency, requiring massive secondary compute infrastructure just to monitor the primary model.⁹
3. **Asynchronous Monitoring:** Deep scrutiny evaluations running in the background. This system is organized like a flowchart; a smaller model (like Claude Haiku) quickly scans

content and triggers a detailed, computationally expensive analysis by a more advanced model (like Claude Sonnet) if suspicious activity is detected.⁹

- 4. **Advanced Security Infrastructure:** The framework mandates multi-party authorization for model weights, binary authorization for endpoints, deception technology (such as honeypots and fake weights to trap intruders), centralized log management via SIEM/SOAR platforms, and physical Technical Surveillance Countermeasures (TSCM) sweeps of company premises.⁹

While these technical architectures undoubtedly represent state-of-the-art corporate cybersecurity, they also represent catastrophic capital expenditure requirements for any emerging competitor.³³ By establishing ASL-3 as the de facto "best practice" for advanced models, Anthropic defines the bare minimum infrastructure required to operate legally and ethically in the frontier space.³² A startup seeking to compete on raw intelligence capabilities must not only spend billions to train a multi-parameter model, but must also immediately finance a defense-in-depth cybersecurity apparatus equivalent to a Fortune 500 financial institution or defense contractor.

The deeply strategic and self-serving nature of the Responsible Scaling Policy became glaringly apparent during the transition from Version 2.0 to Version 3.0 and 3.1 in early 2026. In earlier iterations of the policy, Anthropic committed to a philosophy of "binding itself to the mast," promising a unilateral pause on all artificial intelligence development and deployment if the company could not effectively meet its own internal safety mitigations.³⁴ However, in 2026, Anthropic quietly dropped the hard pause commitment entirely. The company suddenly argued that a unilateral pause would actually be detrimental from a societal perspective, claiming that if they paused, it would allow less responsible competitors to set the pace of innovation, leading to a world that is fundamentally less safe.³⁴

Instead of binding itself, Anthropic reframed its severe mitigations as "industry-wide recommendations".³² By shifting the burden of compliance from an internal operational mandate to an external industry standard, Anthropic brilliantly avoided handicapping its own aggressive product release schedule while simultaneously asserting that the entire ecosystem must be held to its uniquely constructed, highly expensive framework.³⁷ Open-source advocates, policy analysts, and safety researchers immediately identified this pivot as a transition away from genuine risk reduction and toward flexible, narrative-driven market gatekeeping.³¹ The RSP is no longer a set of strict, verifiable commitments; it is an ideological instrument used to signal virtue to federal regulators while continually shifting the infrastructural goalposts for competitors.³⁸ As one reviewer noted, the fundamental design principle of the new RSP is "flexibility," allowing Anthropic to change the contents at any time while forcing competitors to constantly react to Anthropic's definition of safety.³⁸

Evolution of Anthropic's	Core Mechanisms &	Strategic Market Effect
--------------------------	-------------------	-------------------------

Responsible Scaling Policy	Operational Commitments	
Versions 1.0 - 2.0 (2023-2024)	Rigid Capability Thresholds triggering strict ASL-2/ASL-3 mitigations. Included a hard commitment to unilaterally pause development if safety standards could not be met. ⁸	Established "safety" as Anthropic's primary product differentiator against OpenAI, building public and regulatory trust. ³¹
Version 3.0 (February 2026)	Dropped the unilateral pause commitments entirely. Introduced the concept of "industry-wide recommendations" and the publication of periodic Risk Reports. ³⁴	Abandoned self-handicapping while simultaneously pressuring the broader industry to adopt Anthropic's expensive compliance infrastructure. ³⁴
Version 3.1 (April 2026)	Refined the AI R&D threshold (compressing two years of aggregate progress). Clarified absolute corporate flexibility to pause or advance arbitrarily regardless of policy text. ³²	Maintained total operational flexibility for the incumbent while formalizing the expectation of massive infrastructural moats for any market entrants. ³²

Project Glasswing: The Privatization of National Security and the Creation of Elite Cartels

The most explicit and damaging manifestation of artificial intelligence safety being utilized as a mechanism for regulatory capture and market consolidation occurred in April 2026 with the announcement of "Project Glasswing".³⁹ This initiative represents a massive paradigm shift in how foundational models are distributed and highlights the severe competitive consequences of allowing private corporations to govern access to technology based on proprietary threat narratives.

In early April 2026, internal documents leaked to the press revealing that Anthropic had developed an internal frontier model codenamed "Claude Mythos" (specifically the Mythos Preview).⁴⁰ The leaked documents indicated that this model achieved a terrifying capability

discontinuity in the realm of cybersecurity, autonomously discovering thousands of zero-day vulnerabilities—critical software flaws previously unknown and entirely unpatched—across every major operating system, web browser, and virtual machine monitor in existence.³⁹ In evaluations using the rigorous CyberGym benchmark, Mythos Preview scored an unprecedented 83.1% in cybersecurity vulnerability reproduction, vastly outperforming the previous flagship model, Claude Opus 4.6, which scored only 66.6%.³⁹ More alarmingly, the model successfully chained together multiple vulnerabilities to escape secured sandbox environments, demonstrating a dangerous ability to bypass its own safeguards.⁴²

Citing severe, immediate risks to national security, global economies, and public safety, Anthropic invoked its Responsible Scaling Policy and declared that Claude Mythos was simply too dangerous to be released to the general public, the open-source community, or even accessed via standard commercial APIs.⁴⁰ Instead of a general deployment, Anthropic launched Project Glasswing, an initiative purportedly designed to use Mythos purely for defensive cybersecurity purposes to secure the world's most critical digital infrastructure before malicious actors could achieve similar capabilities.³⁹

The deep controversy surrounding Project Glasswing stems not from the discovery of the model itself, but from the extreme exclusivity of its distribution. Upon determining the model was a threat to global security, Anthropic did not surrender the technology to the United States Department of Defense, nor did it create an open, democratic, and transparent consortium for public infrastructure protection. Instead, it restricted access to a hand-selected coalition of twelve founding corporate partners: Amazon Web Services, Apple, Broadcom, Cisco, CrowdStrike, Google, JPMorgan Chase, The Linux Foundation, Microsoft, NVIDIA, and Palo Alto Networks.³⁹ Anthropic committed \$100 million in usage credits exclusively to these entities, alongside a nominal \$4 million donation to open-source security organizations.³⁹

From a strict market dynamics perspective, Project Glasswing effectively functions as a private, un-regulated cartel.⁴⁷ The select organizations granted access to Mythos are simultaneously Anthropic's largest financial investors, major cloud computing providers, and direct enterprise customers.⁴⁷ By granting these specific mega-corporations exclusive access to an advanced, AI-powered vulnerability detection engine, Anthropic bestowed upon them an insurmountable asymmetric advantage in the global cybersecurity marketplace.⁴⁵

For software builders, mid-sized financial institutions, and competing cybersecurity startups that were deliberately excluded from the Glasswing consortium, the competitive implications are fatal.⁴⁸ A competing cybersecurity firm relying on human researchers and traditional fuzzing techniques cannot mathematically compete with a corporate entity utilizing a frontier AI capable of identifying zero-days in hours instead of months.⁴³

The independent security community rapidly highlighted this glaring discrepancy. Analysts noted that there was absolutely "no public process for that selection. No regulatory oversight. No democratic input. A private company made a unilateral decision about who controls" the most powerful cyber-weapon developed to date.⁴⁷ While Anthropic framed the restriction as a

moral necessity of "safety," the harsh market reality is that they fortified the existing Silicon Valley oligopoly.¹² If artificial intelligence vulnerability detection has truly reached a threshold that changes the fundamental economics of defensive security, denying that technology to the broader market while explicitly arming incumbents is the ultimate expression of regulatory capture.⁴⁶ It forces the entire global economy to rely exclusively on the Glasswing partners for critical software security, thereby permanently establishing those twelve specific companies as the unassailable arbiters of digital safety.⁵¹

Project Glasswing Consortium	The Excluded Market	Market Implication
AWS, Google, Microsoft, Apple, NVIDIA, Broadcom, Cisco, CrowdStrike, Palo Alto Networks, JPMorgan Chase. ³⁹	Independent cybersecurity startups, mid-sized enterprises, government agencies not explicitly invited, the broader open-source developer community. ⁴⁵	Creates an insurmountable asymmetric advantage. Consortium members possess automated zero-day detection, while the excluded market relies on legacy human analysis. ⁴⁶
Receive a combined \$100 million in free usage credits for Claude Mythos Preview. ³⁹	Rely on a \$4 million token donation distributed across the entire open-source security ecosystem. ³⁹	Massively subsidizes the operational costs of multi-trillion dollar incumbents while leaving the open-source community chronically underfunded. ³⁹
Operate as Anthropic's investors, cloud providers, and enterprise clients. ⁴⁷	Subject to Anthropic's Responsible Scaling Policy restrictions and API throttling. ¹²	Fuses the supply chain. The entities governing AI safety are identical to the entities profiting from AI deployment. ⁴⁷

The Application Layer Purge: OpenClaw and the Squeeze on Open Source

While structural moats are established through high-level corporate alliances and federal regulatory frameworks, the tactical enforcement of market dominance occurs directly at the application and developer layer. In 2026, the battle for control of the application layer was vividly illustrated by the conflict between closed-enterprise SaaS systems and open-source AI

agents, culminating in the coordinated suppression of the OpenClaw framework.⁵⁴

By early 2026, a project known as "OpenClaw" (formerly Clawdbot) had become the fastest-growing open-source artificial intelligence project on record, amassing over 45,000 GitHub stars and generating a vast repository of community-maintained workflows known as ClawdHub.⁵⁵ OpenClaw functioned as a personal, autonomous agent framework that directly connected large language models to a user's local file system, terminal, web browsers, and over 50 third-party integrations (including messaging applications like WhatsApp and Discord).⁵⁷ Crucially, the tool was model-agnostic, allowing users to plug in any API key to achieve 24/7 automation for complex tasks ranging from automated coding and rigorous code reviews to newsletter drafting, web scraping, and financial market analysis.⁵⁷

The open-source nature of OpenClaw presented a massive, disruptive threat to the monetization strategies of foundational model providers. Because it allowed users to bypass official web interfaces and orchestrate highly complex, looping autonomous workflows locally, it threatened to transform large language models into highly commoditized backend utilities rather than premium enterprise platforms.⁵⁴

In early April 2026, Anthropic executed a devastating policy shift designed to break this ecosystem. Citing the "outsized strain" that autonomous loops placed on its infrastructure, Anthropic officially banned the use of OpenClaw and similar third-party wrappers under its standard Claude flat-rate subscription plans.⁵⁹ Starting precisely at 3:00 PM Eastern Time on April 4, developers running OpenClaw workflows were forced off the subscription tier and pushed entirely onto metered API usage-based billing.⁶¹

Due to the fundamental nature of autonomous agents—which constantly loop, retry tasks, and manage extensive context windows—the sudden transition to per-token API billing caused operational costs for open-source developers to explode exponentially.⁵⁴ Venture capitalist Marc Andreessen highlighted the severity of this dynamic, noting that high-performance autonomous agents like OpenClaw can cost between \$300 and \$1,000 per day to operate at scale under standard API pricing.⁵⁹

This pricing pivot, immediately dubbed the "Claw Tax" by enraged developers, decimated the open-source AI community overnight.⁵⁴ Thousands of independent developers and lean startups lost their flat-rate execution environments, effectively pricing grass-roots innovation out of the market.⁶¹ Anthropic's engineering lead, Boris Cherny, defended the draconian move as a necessary engineering constraint optimization, arguing that Anthropic's systems were highly optimized for direct chat workloads and could not sustain the persistent memory demands of autonomous agents.⁶⁰

However, the narrative of mere engineering constraints collapsed entirely under strategic scrutiny. Barely weeks prior to banning OpenClaw from its subscription tiers, Anthropic launched "Claude CoWork" (featuring the Claude Dispatch agent)—a proprietary, closed-source enterprise productivity platform that mirrored nearly all of OpenClaw's capabilities, including deep software integrations, recurring task management, and file system

interactions.⁵⁴

The juxtaposition was stark and undeniably anti-competitive: Anthropic was systematically suffocating the open-source alternative through punitive API pricing policies while offering the exact same technical utility within a tightly controlled, high-margin enterprise subscription.⁵⁴ Developers quickly noted the extreme convenience of the timing. The creator of OpenClaw, Peter Steinberger (who subsequently joined OpenAI), observed the hypocrisy publicly: "Funny how timings match up, first they copy some popular features into their closed harness, then they lock out open source".⁵⁴

By choking off access to affordable compute for open-source wrappers, frontier labs effectively transition their business models from selling base intelligence to selling vertically integrated, monopolistic enterprise software.⁵⁴ The open-source community views this not as an infrastructure necessity, but as a hostile commercial retaliation disguised as a safety and policy imperative, confirming that the incumbents will not tolerate any layer of abstraction that sits between them and the end-user.⁶¹

Comparative Feature	Claude CoWork (Anthropic Proprietary)	OpenClaw (Open-Source Framework)	Market Impact of Anthropic Policy
Architecture	Closed-source, operates inside the Claude desktop app. ⁶²	Open-source, connects via gateways, terminal, and VPS. ⁶²	Forces users into Anthropic's walled garden UI, eliminating local control.
Execution	Local execution; pauses if application is closed. ⁶²	Runs 24/7 on remote servers via API. ⁶²	Limits true 24/7 automation to those willing to pay exorbitant API costs.
Pricing Model	Flat-rate enterprise subscription (Claude Pro/Max). ⁶²	Requires API tokens, subjected to the "Claw Tax" (up to \$1,000/day). ⁵⁴	Prices out startups and independent developers, ensuring only large enterprises can afford autonomous workflows. ⁶¹

Integrations	20+ enterprise-grade plugins governed by Anthropic. ⁶³	50+ third-party integrations via the community ClawdHub. ⁵⁷	Centralizes the ecosystem, killing the community-driven plugin marketplace.
---------------------	---	--	---

Intra-Oligopoly Warfare and the Institutionalization of Influence

The deployment of regulatory capture strategies is rarely perfectly unified; the oligopoly occasionally fractures as massive firms pursue diverging strategic vectors to undermine one another. In 2026, the ideological and commercial rift between Anthropic and OpenAI broke into public view, shedding critical light on exactly how these companies weaponize policy and public relations against each other.

In April 2026, a highly sensitive internal memo authored by OpenAI's Chief Revenue Officer, Denis D'Almeida (and heavily circulated by executives like Dresser), leaked to the public press.¹⁰ The internal memo outlined OpenAI's aggressive Q2 strategy to counter Anthropic and offered a scathing, insider critique of Anthropic's corporate posture regarding safety.⁶⁶

Within the memo, OpenAI explicitly accused Anthropic of engaging in "anti-competitive gatekeeping dressed up as public safety".¹² The memo suggested that Anthropic's highly publicized decisions to withhold models (like the Claude Mythos Preview) and enforce extreme safety constraints were not born of genuine altruism, but were rather highly calculated bets that enterprise buyers, risk-averse financial institutions, and government regulators would ultimately mandate extreme caution.¹² Furthermore, OpenAI highlighted what it termed a "structural advantage": the memo argued that Anthropic's safety-first posture and continuous throttling of API usage was partially a brilliant public relations smokescreen designed to hide a severe, foundational deficit in raw compute infrastructure.⁵³ According to OpenAI's internal analysis, Anthropic's failure to secure sufficient compute hardware early in the hardware cycle forced them to throttle their APIs and restrict model availability, which they subsequently spun to the public under the noble guise of "responsible scaling".¹²

To counter Anthropic's narrative dominance, OpenAI detailed its own platform war strategy centered on a next-generation reasoning model codenamed "Spud" (widely rumored to be GPT-5.5, trained exclusively on NVIDIA Blackwell architecture) and an enterprise agent platform named "Frontier".¹⁰ Unlike Anthropic's strategy of positioning safety as critical infrastructure, OpenAI opted for a full-stack vertical integration play, aiming to dominate the "SuperApp" ecosystem and completely monopolize the deployment engine layer.¹⁰ The leak of this memo underscored a critical truth: within the upper echelons of Silicon Valley, the public rhetoric of existential risk is widely understood not as a scientific certainty, but as a multifaceted, highly effective competitive weapon.

To fortify its position against both OpenAI's overwhelming raw compute power and the open-source community's disruptive agility, Anthropic radically expanded its institutional influence apparatus in Washington, D.C. Corporate lobbying by artificial intelligence firms exploded exponentially in 2025 and 2026; lobbyist activity on AI issues grew by an astonishing 265%, with software companies deploying thousands of operatives to directly shape federal legislation.⁶⁸

Anthropic aggressively joined this fray when its commercial interests were directly threatened. In early March 2026, the Pentagon designated Anthropic as a "supply chain risk" under 10 U.S.C. § 3252 following highly tense negotiations where Anthropic refused to drop proprietary safety guardrails that restricted mass domestic surveillance and fully autonomous weapons systems.⁶⁹ Faced with the potentially crippling loss of lucrative federal defense contracts, Anthropic immediately hired Ballard Partners—a powerful Washington lobbying firm with exceptionally close ties to the Trump administration—to navigate the dispute and pressure the Pentagon.⁶⁹ Over the course of the year, Anthropic's federal lobbying spend skyrocketed by over 330% year-over-year, exceeding \$3.1 million distributed across six different lobbying firms.⁶⁹

Simultaneously, recognizing the need to shape the intellectual foundation of future regulations, Anthropic announced the creation of the "Anthropic Institute," a dedicated policy think tank based in Washington, D.C..⁴ Directed by Anthropic co-founder Jack Clark (who strategically transitioned from head of public policy to the newly created role of head of public benefit), the Institute is designed to directly shape the national dialogue on job displacement, economic stability, and democratic AI governance.⁷¹ The stated philanthropic goal of the Institute is to provide crucial information to guide society through the inevitable transition to powerful AI systems.⁷²

However, policy experts and political scientists recognize the establishment of corporate think tanks by AI companies (OpenAI similarly acquired a tech podcast and opened an advocacy workspace in D.C.) as an aggressive, highly orchestrated effort to reshape the narrative amid growing public disapproval of data center energy consumption, grid strain, and AI-driven job losses.⁴ By fully funding its own independent-sounding institute, Anthropic ensures that the federal policymakers drafting the preemption frameworks and defining the compute thresholds are educated directly by the company's own philosophical architects. It is the ultimate, seamless fusion of corporate public relations, safety philosophy, and hard-nosed legislative lobbying—guaranteeing that any future regulatory body will view the entire artificial intelligence industry exclusively through Anthropic's highly specific, incumbent-friendly lens.

Strategic Conclusions and Future Trajectories

The comprehensive analysis of the artificial intelligence ecosystem in 2026 confirms that the narrative of "AI safety" has definitively transcended its origins as a valid technical alignment problem to become the single most dominant strategy for corporate consolidation and market gatekeeping in the modern technology sector. The mechanisms of this regulatory capture are

highly diverse, operating simultaneously across legislative, infrastructural, and economic domains to form an impenetrable barrier to entry.

The interplay between the chaotic state-level legislative patchwork and the White House National Policy Framework perfectly illustrates how incumbents generate market conditions that absolutely necessitate federal intervention.⁶ By strongly supporting federal preemption, technology oligopolies ensure the total neutralization of localized, democratically enacted consumer protection laws while actively and quietly designing the "minimally burdensome" national standards that replace them.²⁵ This legislative maneuvering ensures that the only viable legal environment for artificial intelligence development is one calibrated precisely to the operational capacities and massive legal budgets of a multi-billion-dollar corporation.

Similarly, the evolution of Responsible Scaling Policies demonstrates the dangerous privatization of technological regulation. By pivoting away from unilateral pause commitments and moving toward industry-wide architectural mandates (such as the ASL-3 defense-in-depth infrastructure requirements), frontier developers successfully force lean startups to internalize massive cybersecurity capital expenditures before they can even begin to compete on core model capabilities.⁸

The deployment of Project Glasswing stands as the most acute and alarming manifestation of this gatekeeping behavior. The unilateral corporate decision to restrict a world-altering zero-day vulnerability discovery engine (Claude Mythos) to a private cartel of twelve foundational tech companies ensures that the cybersecurity paradigm of the future remains the exclusive, highly profitable domain of the present oligopoly.³⁹ When viewed alongside the aggressive strangulation of the OpenClaw open-source ecosystem through punitive API pricing and the parallel, highly coordinated launch of proprietary enterprise clones like Claude CoWork, the holistic strategy becomes undeniable.⁵⁴

For federal policymakers, antitrust regulators, and the broader technology sector, the events of 2025 and 2026 demand a radical, immediate reassessment of artificial intelligence governance. The current trajectory heavily suggests that without aggressive intervention targeted specifically at anti-competitive behavior disguised as public safety, the artificial intelligence industry will crystallize into a permanent, heavily fortified oligopoly. The frontier models of the future will not be constrained by the limits of compute availability, power generation, or algorithmic ingenuity, but by the legally enforced moats built today under the unimpeachable, meticulously crafted banner of protecting humanity.

Works cited

1. 2025 AI Safety Index - Future of Life Institute, accessed April 14, 2026, <https://futureoflife.org/ai-safety-index-summer-2025/>
2. Research - Anthropic, accessed April 14, 2026, <https://www.anthropic.com/research>
3. AI Regulation is Coming- What is the Likely Outcome? - CSIS, accessed April 14, 2026,

<https://www.csis.org/blogs/strategic-technologies-blog/ai-regulation-coming-what-likely-outcome>

4. AI companies know they have an image problem. Will funding policy papers and thinktanks dig them out?, accessed April 14, 2026,
<https://www.theguardian.com/technology/2026/apr/12/ai-image-problem-policy-papers-thinktanks>
5. Artificial Intelligence - Advocacy - California Chamber of Commerce, accessed April 14, 2026,
<https://advocacy.calchamber.com/policy/issues/artificial-intelligence/>
6. The AI Regulation Nightmare: Why State-by-State Rules are Crushing Your Business Innovation - Kelley Kronenberg, accessed April 14, 2026,
<https://kelleykronenberg.com/the-ai-regulation-nightmare-why-state-by-state-rules-are-crushing-your-business-innovation/>
7. SB 53: What California's New AI Safety Law Means for Developers, accessed April 14, 2026,
<https://ai-analytics.wharton.upenn.edu/wharton-accountable-ai-lab/sb-53-what-californias-new-ai-safety-law-means-for-developers/>
8. Measurement Challenges in AI Catastrophic Risk Governance and Safety Frameworks, accessed April 14, 2026,
<https://www.techpolicy.press/measurement-challenges-in-ai-catastrophic-risk-governance-and-safety-frameworks/>
9. Anthropic's Responsible Scaling Policy (version 3.0), accessed April 14, 2026,
<https://www.anthropic.com/responsible-scaling-policy/rsp-v3-0>
10. OpenAI's \$852 billion valuation is under scrutiny from its own investors - TNW, accessed April 14, 2026,
<https://thenextweb.com/news/openai-852-billion-valuation-investor-scrutiny-anthropic-revenue>
11. Marc Andreessen's 2026 Outlook: AI Timelines, US vs. China, and The Price of AI, accessed April 14, 2026,
<https://podcasts.apple.com/ae/podcast/marc-andreessens-2026-outlook-ai-timelines-us-vs-china/id842818711?i=1000744113588>
12. Anthropic Withholds AI Model Citing Safety, But Skeptics See Strategy, accessed April 14, 2026,
<https://startupfortune.com/anthropic-withholds-ai-model-citing-safety-but-skeptics-see-strategy/>
13. AI News Briefs BULLETIN BOARD for April 2026 | Radical Data Science, accessed April 14, 2026,
<https://radicaldatascience.wordpress.com/2026/04/03/ai-news-briefs-bulletin-board-for-april-2026/>
14. State of AI governance regulations in the United States: what you need to know in 2026, accessed April 14, 2026,
<https://verifywise.ai/blog/state-of-ai-governance-regulations-united-states-2026>
15. Safe and Secure Innovation for Frontier Artificial Intelligence Models Act - Wikipedia, accessed April 14, 2026,
https://en.wikipedia.org/wiki/Safe_and_Secure_Innovation_for_Frontier_Artificial_I

[ntelligence_Models_Act](#)

16. California Tech Legislation Roundup: Numerous Privacy and AI ..., accessed April 14, 2026,
<https://epic.org/california-tech-legislation-roundup-numerous-privacy-and-ai-laws-enacted-and-six-vetoed/>
17. California's Transparency in Frontier Artificial Intelligence Act: 5 Things You Should Know, accessed April 14, 2026,
<https://www.orrick.com/en/Insights/2025/10/Californias-Transparency-in-Frontier-Artificial-Intelligence-Act-5-Things-You-Should-Know>
18. California enacts landmark AI transparency law: The Transparency in Frontier Artificial Intelligence Act | White & Case LLP, accessed April 14, 2026,
<https://www.whitecase.com/insight-alert/california-enacts-landmark-ai-transparency-law-transparency-frontier-artificial>
19. Bill Text: CA SB53 | 2025-2026 | Regular Session | Enrolled - LegiScan, accessed April 14, 2026, <https://legiscan.com/CA/text/SB53/id/3270002>
20. Opportunity Costs of State and Local AI Regulation | Cato Institute, accessed April 14, 2026,
<https://www.cato.org/policy-analysis/opportunity-costs-state-local-ai-regulation>
21. California Passes First-of-Its-Kind AI Safety Law | MultiState, accessed April 14, 2026,
<https://www.multistate.us/insider/2025/10/9/california-passes-first-of-its-kind-ai-safety-law>
22. Ensuring a National Policy Framework for Artificial Intelligence - The White House, accessed April 14, 2026,
<https://www.whitehouse.gov/presidential-actions/2025/12/eliminating-state-law-obstruction-of-national-artificial-intelligence-policy/>
23. The White House Legislative Recommendations: National Policy Framework for Artificial Intelligence and Federal Preemption of State AI Laws | Insights | Ropes & Gray LLP, accessed April 14, 2026,
<https://www.ropesgray.com/en/insights/alerts/2026/03/the-white-house-legislative-recommendations-national-policy-framework-for-artificial-intelligence-and-federal-preemption-of-state-ai-laws>
24. White House Releases a National Policy Framework for Artificial Intelligence | Insights, accessed April 14, 2026,
<https://www.hklaw.com/en/insights/publications/2026/03/white-house-releases-a-national-policy-framework-for-artificial-intelligence>
25. White House National Policy Framework for AI ... - The White House, accessed April 14, 2026,
<https://www.whitehouse.gov/wp-content/uploads/2026/03/03.20.26-National-Policy-Framework-for-Artificial-Intelligence-Legislative-Recommendations.pdf>
26. The White House's National Policy Framework for Artificial Intelligence: what it means and what comes next | Consumer Finance Monitor, accessed April 14, 2026,
<https://www.consumerfinancemonitor.com/2026/04/08/the-white-houses-national-policy-framework-for-artificial-intelligence-what-it-means-and-what-comes-next/>

27. White House AI Framework Puts Federal Preemption at the Center of the Debate, accessed April 14, 2026, <https://www.morganlewis.com/pubs/2026/03/white-house-ai-framework-puts-federal-preemption-at-the-center-of-the-debate>
28. White House AI Framework Proposes Industry-Friendly Legislation - Lawfare, accessed April 14, 2026, <https://www.lawfaremedia.org/article/white-house-ai-framework-proposes-industry-friendly-legislation>
29. On the White House AI Framework (If You Could Call It That) - Public Knowledge, accessed April 14, 2026, <https://publicknowledge.org/on-the-white-house-ai-framework-if-you-could-call-it-that/>
30. Beware the AI Preemption Trap - Just Security, accessed April 14, 2026, <https://www.justsecurity.org/134693/beware-ai-preemption-trap/>
31. Responsible Scaling Policy Version 3.0 - Anthropic, accessed April 14, 2026, <https://www.anthropic.com/news/responsible-scaling-policy-v3>
32. Responsible Scaling Policy Updates \ Anthropic, accessed April 14, 2026, <https://www.anthropic.com/responsible-scaling-policy>
33. Surviving AI Price Wars Without Destroying Your Business | Andreessen Horowitz, accessed April 14, 2026, <https://a16z.com/surviving-ai-price-wars-without-destroying-your-business/>
34. Anthropic's RSP v3.0: How it Works, What's Changed, and Some Reflections | GovAI, accessed April 14, 2026, <https://www.governance.ai/analysis/anthropics-rsp-v3-0-how-it-works-whats-changed-and-some-reflections>
35. Responsible Scaling Policy v3 - LessWrong, accessed April 14, 2026, <https://www.lesswrong.com/posts/HzKuzrKfaDJvQqmjh/responsible-scaling-policy-v3>
36. Responsible Scaling Policy v3 — EA Forum, accessed April 14, 2026, <https://forum.effectivealtruism.org/posts/DGZNAGL2FNJfftwtgE/responsible-scaling-policy-v3-1>
37. Responsible Scaling Policy | Version 3.0 - Anthropic, accessed April 14, 2026, <https://anthropic.com/responsible-scaling-policy/rsp-v3-0>
38. Anthropic Responsible Scaling Policy v3: Dive Into The Details - LessWrong, accessed April 14, 2026, <https://www.lesswrong.com/posts/RtQxa5MoKk9bwEEEd/anthropic-responsible-scaling-policy-v3-dive-into-the>
39. Project Glasswing: Securing critical software for the AI era - Anthropic, accessed April 14, 2026, <https://www.anthropic.com/glasswing>
40. What Is Claude Mythos? Anthropic's Most Dangerous AI Model Explained - MindStudio, accessed April 14, 2026, <https://www.mindstudio.ai/blog/what-is-claude-mythos-anthropic-most-dangerous-ai-model>
41. Why Anthropic's new model has cybersecurity experts rattled - Platformer, accessed April 14, 2026,

- <https://www.platformer.news/anthropic-mythos-cybersecurity-risk-experts/>
42. Anthropic's Claude Mythos Finds Thousands of Zero-Day Flaws Across Major Systems - The Hacker News, accessed April 14, 2026, <https://thehackernews.com/2026/04/anthropics-claude-mythos-finds.html>
 43. Claude Mythos Preview \ red.anthropic.com, accessed April 14, 2026, <https://red.anthropic.com/2026/mythos-preview/>
 44. Project Glasswing powered by Claude Mythos: defending software before hackers do, accessed April 14, 2026, <https://securityaffairs.com/190496/ai/project-glasswing-powered-by-claude-mythos-defending-software-before-hackers-do.html>
 45. Project Glasswing Just Made Your Security Playbook Obsolete - GovInfoSecurity, accessed April 14, 2026, <https://www.govinfosecurity.com/project-glasswing-just-made-your-security-playbook-obsolete-a-31396>
 46. Anthropic Glasswing: AI Vulnerability Detection Has Crossed a Threshold, accessed April 14, 2026, <https://futurumgroup.com/insights/anthropic-glasswing-ai-vulnerability-detection-has-crossed-a-threshold/>
 47. Glasswing Is the Confirmation: The 'Manhattan Project' for AI Arrived on April 7, 2026 | by Berend Watchus - OSINT Team, accessed April 14, 2026, <https://osintteam.blog/glasswing-is-the-confirmation-the-manhattan-project-for-ai-arrived-on-april-7-2026-16f2edbb4ddb>
 48. AI Security Threats: Project Glasswing and Mythos | Black Duck Blog, accessed April 14, 2026, <https://www.blackduck.com/blog/ai-security-glasswing-mythos.html>
 49. Project Glasswing: Securing critical software for the AI era | Hacker News, accessed April 14, 2026, <https://news.ycombinator.com/item?id=47679121>
 50. Project Glasswing : Anthropic's Security Initiative, Claude Mythos Preview, Key Benchmarks, Partners, Pricing, and What It Means for Enterprise Defense | ALM Corp, accessed April 14, 2026, <https://almcorp.com/blog/project-glasswing-anthropics-security-initiative-claude-mythos-preview-key-benchmarks-partners-pricing-and-what-it-means-for-enterprise-defense/>
 51. Project Glasswing and the Case for a Diverse Agentic AI Strategy | Alchemy Tech Group, accessed April 14, 2026, <https://alchemytechgroup.com/blog/project-glasswing>
 52. Project Glasswing Has Big Tech United on Security. One Name Is Missing. | Blog of Daniel Ryan Reiff — founder, designer, and developer, accessed April 14, 2026, <https://danielreiff.com/blog/project-glasswing-has-big-tech-united-on-security-one-name-is-missing>
 53. OpenAI leaked memo slams Anthropic: ChatGPT maker accuses rival of 'fear-based' AI approach, reveals report | Mint, accessed April 14, 2026, <https://www.livemint.com/technology/tech-news/openai-leaked-memo-slams-anthropic-chatgpt-maker-accuses-rival-of-fear-based-ai-approach-reveals-report-11776136226214.html>

54. Anthropic Suspended the OpenClaw Creator's Claude Account , And It Reveals a Much Bigger Problem : r/Al_Agents - Reddit, accessed April 14, 2026, https://www.reddit.com/r/Al_Agents/comments/1sjqxt/anthropic_suspended_the_openclaw_creators_claude/
55. OpenClaw - Wikipedia, accessed April 14, 2026, <https://en.wikipedia.org/wiki/OpenClaw>
56. Your employees are already using this. IT doesn't know yet., accessed April 14, 2026, <https://gradientflow.substack.com/p/your-employees-are-already-using>
57. OpenClaw Use Cases 2026: 25+ Real Examples (Updated February) - TLDL, accessed April 14, 2026, <https://www.tldl.io/blog/openclaw-use-cases-2026>
58. What is OpenClaw? Your Open-Source AI Assistant for 2026 | DigitalOcean, accessed April 14, 2026, <https://www.digitalocean.com/resources/articles/what-is-openclaw>
59. Marc Andreessen sketches AGI pricing for advanced AI systems like OpenClaw, accessed April 14, 2026, <https://m.economictimes.com/tech/artificial-intelligence/marc-andreessen-sketches-agi-pricing-for-advanced-ai-systems-like-openclaw/articleshow/130119540.cms>
60. Anthropic cuts OpenClaw access from Claude subscriptions, offers credits to ease transition, accessed April 14, 2026, <https://www.infoworld.com/article/4154435/anthropic-cuts-openclaw-access-from-claude-subscriptions-offers-credits-to-ease-transition.html>
61. Overnight Change, Anthropic Officially Bans OpenClaw! Global Developers Collapse in 24 Hours | ME News on Binance Square, accessed April 14, 2026, <https://www.binance.com/en/square/post/308749699730466>
62. Claude CoWork vs OpenClaw: The Brutally Honest Automation Breakdown - Reddit, accessed April 14, 2026, https://www.reddit.com/r/AISEOInsider/comments/1rifamf/claude_cowork_vs_openclaw_the_brutally_honest/
63. Claude Cowork vs. OpenClaw: Anthropic's Enterprise Power Move You Can't Ignore, accessed April 14, 2026, <https://medium.com/@shahadilh18/claude-cowork-vs-openclaw-anthropics-enterprise-power-move-you-can-t-ignore-b087ba53c281>
64. Anthropic's response to the AI tool that caused lines around the block in Shenzhen, accessed April 14, 2026, <https://thenewstack.io/claude-dispatch-versus-openclaw/>
65. Anthropic did the absolute right thing by sending OpenClaw a cease & desist and allowing Sam Altman to hire the developer - Reddit, accessed April 14, 2026, https://www.reddit.com/r/ClaudeAI/comments/1r99wg1/anthropic_did_the_absolute_right_thing_by_sending/
66. OpenAI Internal Memo Leaks: Project Spud Model Competes with Mythos in Direct Counter to Anthropic, accessed April 14, 2026, <https://news.aibase.com/news/27108>
67. Anthropic vs OpenAI vs Google: Three Different Bets on the Future of AI Agents | MindStudio, accessed April 14, 2026,

- <https://www.mindstudio.ai/blog/anthropic-vs-openai-vs-google-agent-strategy>
68. Generative Influence - Public Citizen, accessed April 14, 2026,
<https://www.citizen.org/article/generative-influence/>
69. Anthropic hires Trump-linked lobbyist after Pentagon risk label, accessed April 14, 2026,
<https://www.techinasia.com/news/anthropic-hires-trumplinked-lobbyist-pentagon-risk-label>
70. Anthropic Hires Ballard Partners for Lobbying Amid Pentagon AI Dispute, accessed April 14, 2026,
<https://www.kucoin.com/news/flash/anthropic-hires-ballard-partners-for-lobbying-amid-pentagon-ai-dispute>
71. Anthropic's New AI Think Tank & Leadership Changes - Written by AI, accessed April 14, 2026,
<https://www.writtenby.ai/blog/2026/03/11/anthropic-ai-think-tank-leadership-changes/>
72. Introducing The Anthropic Institute, accessed April 14, 2026,
<https://www.anthropic.com/news/the-anthropic-institute>
73. Reassurance calls: 'Stop hiring humans'? Silicon Valley confronts AI job panic - RTL Today, accessed April 14, 2026,
<https://today.rtl.lu/news/world/stop-hiring-humans-silicon-valley-confronts-ai-job-panic-1337336447>