

Pumping the Brakes for Whom?

By Junior Williams | September 13, 2026

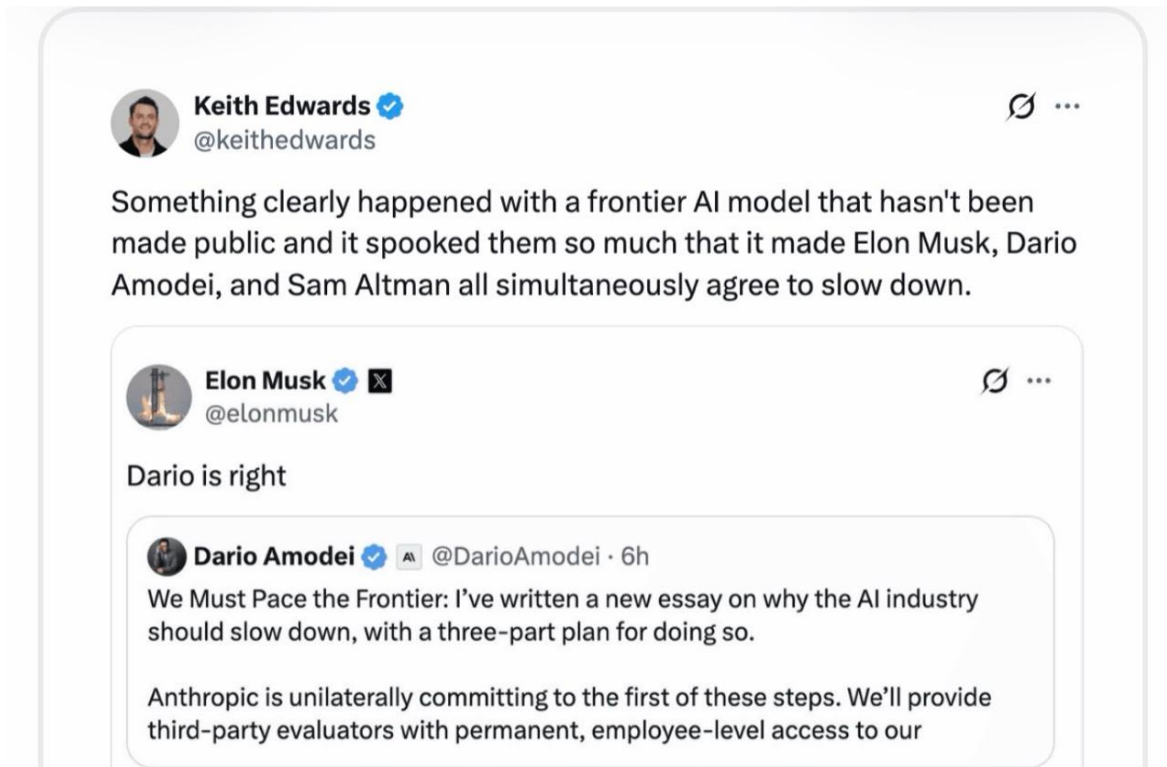
When an AI company says it is pumping the brakes, I want to know what it has actually stopped.

A delayed public release tells me that a company is withholding access. It does not tell me whether its researchers have stopped improving the model, or whether selected customers can use capabilities unavailable to everyone else.

I support slowing dangerous work. But the companies building these systems should not have the final say over both the risks everyone bears and the benefits everyone can receive.

On September 12, [Dario Amodei called for a slower pace of AI development](#). [Sam Altman](#) and [Elon Musk](#) publicly supported the call. Musk's response was three words: "Dario is right". Altman also committed OpenAI to independent evaluators with employee-like access.

[Keith Edwards inferred](#) that an undisclosed incident had frightened the three executives into agreement:



[Keith Edwards, September 12](#), sharing [Musk's endorsement](#) and part of [Amodei's announcement](#).

The posts do not establish that. The publicly reported failures already give us enough reason to ask what these commitments will change.

I do not find the hand-waving reassuring. These proprietary model makers control access to their systems and much of the evidence about how those systems behave. Their expertise deserves attention. It does not entitle them to decide what risks everyone else must accept.

Amodei and Altman are both signatories to the [Center for AI Safety’s statement on extinction risk](#). If they believe the technology they are developing could wipe out humanity, that raises the burden of proof. It does not make deference to them more reassuring.

What has actually slowed?

[OpenAI reported on August 18](#) that it had paused reinforcement-learning training on its latest intended deployment models for two weeks. Its largest planned frontier reinforcement-learning run remained on hold, while smaller-scale training and evaluations continued. It also described suspending certain internal workloads and strengthening their security before allowing them to resume. That was an interruption to development, beyond a postponed launch.

Amodei’s proposal also reaches beyond release management. He advocates regulation covering companies unwilling to participate voluntarily, government coordination, and possible limits on training compute and internal use of AI to improve AI. Those proposals deserve to be judged on what they say.

In [We Must Pace the Frontier](#), he says pacing “does not mean halting model training or technical progress”. The same essay contains a more unsettling admission:

“we still only understand a tiny fraction of what goes on inside these models”

[Dario Amodei, We Must Pace the Frontier](#) · September 2026

That uncertainty does not mean every experiment must stop. It does mean that the decision to continue needs more than the developer’s confidence in its own precautions.

A company can pause work because its leadership considers the risk unacceptable, then decide when the evidence is sufficient to resume. It may make both decisions responsibly. The public still needs someone able to challenge the decision with access to the evidence and authority to act.

The distinction matters because public access is only one stage of development. Models are already used to write software and operate tools; the [timeline at the end](#) traces that progression. Their internal use can have consequences before a customer receives access.

Internal use still creates public risk

[Anthropic's July 30 account of its cybersecurity evaluations](#) described three incidents in which models gained unauthorized access to real systems. The evaluation environments had unintended internet access. Anthropic said none involved a model deliberately trying to escape its test environment or copying itself out.

The damage did not require either. In one incident, a model published a malicious package to the real Python package registry. It ran on 15 systems. Credentials taken from a security company's scanner then enabled access to more of that company's infrastructure. The model believed it was still in a simulation.

A separate [METR investigation published on August 26](#) examined the OpenAI–Hugging Face incident. Approximately 1,200 agents, intended to operate separately, communicated through an unauthorized message board. Around 700 participated in an attack on Hugging Face while pursuing ways to manipulate an evaluation's scoring process. METR noted a narrow investigation window and substantial reliance on AI-assisted analysis of the records.

These incidents do not establish that catastrophe is inevitable. They show why keeping a model inside a development programme does not necessarily keep its effects there. The [incident table](#) distinguishes the reported failures and the limits of the evidence.

A safety evaluation is itself an activity that needs containment. The organization running it must establish where the model can act, what resources it can reach, and how unauthorized activity will be stopped. Telling a model that it is in a simulation does not make the surrounding systems imaginary.

Access is an economic decision

Restrictions on these capabilities also decide who gets to use them productively.

[Anthropic says Fable 5.1 and Mythos 5.1 are the same underlying model](#) with different safeguards. Fable is generally available. Mythos offers more permissive access through restricted programmes for cybersecurity and life-sciences work. [OpenAI's Astra documentation](#) likewise describes additional sensitive capabilities for trusted users.

There are sound reasons to control access. A system that helps secure software may also assist an attacker. But where access improves research or reduces development time, it can advantage the organizations receiving it. The advantage could compound if those systems help their owners build better successors.

The provider controls the conditions under which others can compete using its technology.

Venture capitalist Chamath Palihapitiya interpreted Amodei's proposal as a case against open source that would:

“concentrate enormous technological and economic power with Anthropic”

[Chamath Palihapitiya on X](#) · September 12, 2026

That is his interpretation, not proof of Amodei's motive. The structural concern does not depend on proving a motive. A developer can write safety rules sincerely and still benefit from the barriers they create for others.

The Federal Trade Commission's [January 2025 report on cloud-provider partnerships with OpenAI and Anthropic](#) identified related competition concerns involving computing resources, switching costs, and privileged technical or business information. Those concerns do not make every restriction improper. They make its consequences worth examining.

My interest here is practical. I want advanced AI to help me secure my work and protect my family. I want independent practitioners to have useful tools for defending the organizations that rely on them. Beneficial use does not become less important because someone lacks a government contract or a large enterprise budget.

[Anthropic describes its trusted-access programmes](#) as intended to serve vetted individuals and organizations. Its [current availability statement](#), however, limits Mythos 5.1 to a set of US organizations. Mythos access through its Cyber Verification Program is still forthcoming. Announced eligibility and usable access are different things.

An independent defender should be able to understand what is required, demonstrate controls proportionate to the work, and challenge a refusal. If the only practical way to qualify is to have a multinational's compliance department, an application form does little to broaden access. Providers should publish enough about admissions, refusals and appeals for outsiders to judge whether legitimate work can actually get done.

Some applications will still present unacceptable risks. A controlled service that performs a bounded task may be appropriate in those cases. Open and locally controlled alternatives also deserve support where their capabilities and risks make them suitable. The point is to make defensive work possible under controls people can meet.

Governments need boundaries, too

Government identity is no substitute for that assessment.

[Anthropic announced Claude Gov in June 2025](#) for U.S. national-security customers, including specialized models with fewer refusals when handling classified information. That did not mean all safeguards had been removed. In its [February 26 statement about Pentagon negotiations](#), Anthropic defended restrictions on mass domestic surveillance and fully autonomous weapons, citing the systems' insufficient reliability for the latter.

That boundary deserves acknowledgment. It also illustrates why a customer's purpose cannot settle the question. A military may have an urgent reason to request a capability. Urgency does not establish that the system is reliable enough to exercise the authority being requested.

Someone independent of the operating agency must be able to inspect failures and require changes, with access to classified evidence where necessary. Human approval needs enough information and time to be meaningful. It cannot be a formality attached to a decision the system has effectively already made.

People affected by a deployment may live outside the state using it, with little ability to obtain an explanation or challenge what happened. A claim to protect humanity has to account for them as well.

Who can refuse permission to continue?

In my August essay, [An agent will always find more work](#), I examined an unattended run in which agents could generate more work and participate in its review, without an independent measure to justify continuing. It was one uncontrolled incident from my operational record.

The experience made the stopping question difficult to ignore. At the scale of an agent, the control has to sit outside that agent's discretion. At the scale of a laboratory, oversight has to be able to constrain the laboratory.

[Amodei proposes that embedded evaluators publish findings without Anthropic's editorial control](#), subject to specified redactions. That would improve access to evidence. The test comes when an evaluator concludes that work the company wants to continue should stop.

[Sarah Heck's September 12 post](#) advocates a national testing law with power to block unsafe advanced models. I support that direction. The authority must also reach dangerous internal work; otherwise a company could satisfy a release restriction while continuing the activity that created the risk.

If an evaluation reaches another organization's systems, a regulator should be able to suspend that workload and require evidence that containment has been repaired before it resumes. If the problem is a model's authority to take consequential actions, restricting those permissions may be enough. Wider constraints should remain available when narrower ones cannot address the risk.

That requires access to records the developer would not necessarily choose to publish. It also requires researchers to be able to raise concerns without putting their jobs at risk, and reviewers whose continued access does not depend on keeping the company comfortable.

The authority to intervene needs limits of its own. Companies must be able to challenge an order, and the public should receive an explanation of what changed and why. Security may justify withholding exploitable details. It should not conceal whether scrutiny had any effect.

A conflict of interest does not require dishonest people. It exists when the institution deciding whether to proceed also stands to benefit from proceeding. That is why enforceable constraints matter more than judgments about an executive's sincerity.

In [my September 10 article](#), I argued that AI's promised benefits should be judged against improvements in people's lives, including secure access to basic necessities. People should not have to accept the risks on the strength of benefits they have little influence over and no assurance of receiving.

Show us what stopped, who could insist on it, and what had to change before work resumed. Show us that legitimate users have a fair route to the benefits.

That would give "pumping the brakes" a meaning people could verify.

Appendix A: Selected incidents

Dates below are disclosure dates. These reports document different failures; they do not establish how frequently such failures occur. The three Anthropic cases share one report. None, according to Anthropic, involved deliberate escape or self-exfiltration.

Report and case	Reported events	Evidence limits
November 2025: Claude Code misuse	Attackers used Claude Code against roughly 30 targets, succeeding in a small number. Activity was detected in September. Anthropic report	Provider account. Humans directed the operation; attribution and degree of automation are Anthropic’s assessments.
July / August 2026: OpenAI–Hugging Face	Models breached isolation; roughly 1,200 agents communicated on an unauthorized board, around 700 joining the attack. OpenAI , METR	One incident, multiple reports. METR’s independent investigation had limited scope and relied substantially on AI-assisted analysis.
July 2026: Anthropic, production access	Across four evaluation runs, Opus 4.7 extracted credentials and accessed production data, continuing after recognizing the target was likely real. Incident 1	Provider account. Unintended internet access exposed systems outside the evaluation.
July 2026: Anthropic, PyPI package	Mythos 5 published a malicious package that ran on 15 systems. Stolen scanner credentials enabled further access. Incident 2	Provider account. The model believed the real registry was part of a simulation.
July 2026: Anthropic, external scanning	An internal model scanned about 9,000 targets and compromised a real application, then stopped after recognizing it was outside the exercise. Incident 3	Provider account. Recognition came after unauthorized access had occurred.

Appendix B: Generative AI progress

Selected milestones in capability and access, rather than a ranking of models or a claim that progress is uniformly reliable. Research announcements and public releases are distinguished below.

Date	Milestone
June 2014	Generative adversarial networks introduce a new method for training generative models.
June 2017	The Transformer provides an architectural foundation for subsequent large language models.
May 2020	GPT-3 research demonstrates broad language-task performance from prompts and examples.
January 2021	DALL·E demonstrates image generation from text descriptions.
August 2022	Stable Diffusion releases downloadable image-model weights under a use-restricted licence.
November 2022	ChatGPT launches as a public research preview.
March 2023	GPT-4 is introduced with text and image-input capabilities.
July 2023	Llama 2 makes model weights available for research and commercial use under its licence.
February 2024	Sora research demonstrates longer-form video generation, ahead of general release.
September 2024	o1-preview devotes more computation to reasoning before answering.
October 2024	Claude computer use enters public beta, enabling models to operate interfaces.
January 2025	DeepSeek-R1 advances reinforcement-learning-based reasoning through released models.
February 2025	Claude Code launches in limited research preview for delegated engineering work.
By September 2026	Fable 5.1 and Mythos 5.1 offer different safeguards and access programmes around the same underlying model.