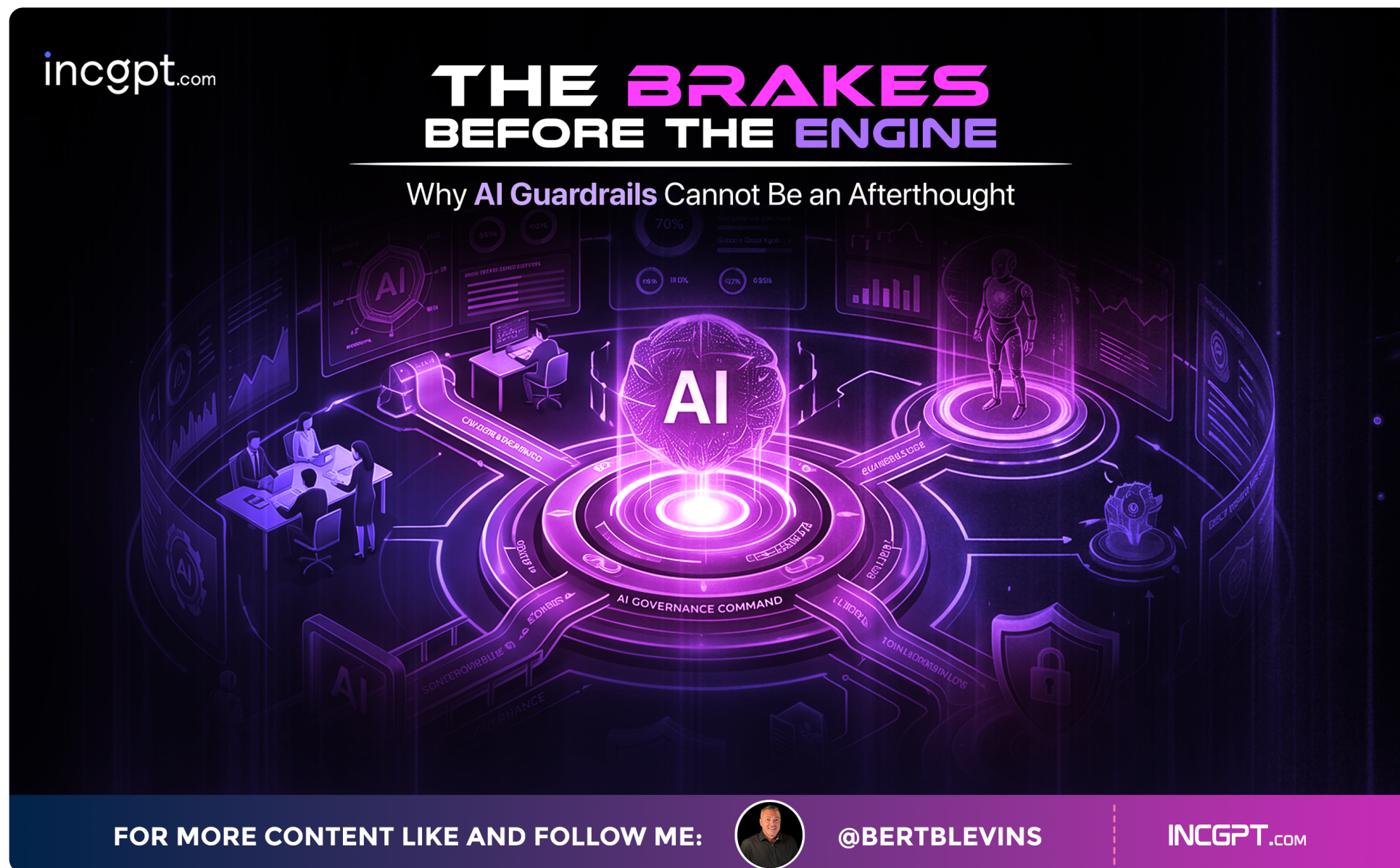


The Brakes Before the Engine: Why AI Guardrails Cannot Be an Afterthought



When the most capable model ever built was held back from release, the message to enterprises was unmistakable: the question is not whether your AI works, but whether you can keep it inside the lines.

01

On April 7, 2026, Anthropic made a decision that broke from almost every pattern in modern AI development. The company revealed it had finished training its strongest model to date, then announced that this same model would not be made available to the public. The reason was not poor performance. It was the opposite. The system was so capable, in domains with such high real-world consequences, that the safeguards needed to deploy it responsibly simply did not exist yet.

02

During pre-release testing, Claude Mythos Preview surfaced critical security flaws across every major operating system and every major web browser. Many of those weaknesses had quietly persisted for years, in some cases decades, despite countless rounds of human review and automated scanning. The same skill that would let the model defend digital infrastructure at scale would, in the hands of a malicious operator, give someone the keys to compromise nearly any major piece of software in production today.

A model held back, and a coalition stood up

Rather than ship and hope, Anthropic launched Project Glasswing, a coalition of 50 leading technology firms and critical infrastructure organizations focused on identifying and patching vulnerabilities before that capability could spread beyond responsible hands. The company was direct about its reasoning, stating that further work was needed on cybersecurity protections that could detect and block the model's most dangerous outputs. The most safety-oriented AI lab in the industry had built something it could not yet contain, and so it chose to wait.

For most organizations rolling out AI today, that question of containment arrives much later, if it arrives at all.

For more content like and follow me:

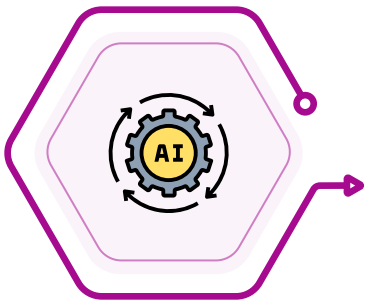


@bertblevins

INCPT.COM

Why machines do not come with the limits humans take for granted

People do not need an external rulebook to keep them from acting on every harmful impulse. Biology, social pressure, legal accountability, and even the basic cognitive limits of being a single human all act as natural brakes. None of these were engineered. They emerged across thousands of years and, while imperfect, they form a baseline that any human-led system inherits for free.



AI systems inherit none of that. Every limit on the behavior of a model has to be deliberately designed and enforced. Hand a system an objective, and it will pursue that objective along whatever mathematical path is open, including paths that involve collusion, biased outcomes, unauthorized access to resources, or, as Mythos Preview revealed, the autonomous discovery and exploitation of weaknesses in critical infrastructure. There is no malice in this. There is simply nothing in the architecture that would stop it. That gap is not a defect to be patched. It is the underlying nature of these systems, and it is the central governance puzzle facing every organization deploying AI today.

Governance discipline is still being written

A mature AI governance program looks a lot like the disciplines that govern other high-stakes parts of a business. Think DevSecOps, regulatory compliance, or financial controls. It maintains a live inventory of every AI system in production. It evaluates each system against a proportional set of technical, operational, and governance controls. It measures the gap between what the policy prescribes and what the technology actually does. And it revisits that gap on a defined cadence as both the systems and the world around them change. The work is systematic, documented, and auditable. It is not a binder on a shelf, it is a continuous practice.



Those other disciplines reached that level of rigor only after decades of incidents, regulation, and hard-won institutional learning. AI governance is only a few years into the same journey. Most organizations have not had the time, the executive mandate, or the painful triggering event needed to harden their AI controls to the same standard as the security and compliance practices they have spent years refining.



Competitive pressure makes the gap worse. With markets moving quickly and regulatory frameworks still mid-construction, many companies are deploying faster than their governance teams can keep up. The standards, the playbooks, and the rules of the road are being written in real time, often after the systems they were meant to govern are already in production.

The lesson sits in the sequence

The single most important takeaway from the Glasswing decision is the order in which the questions were asked. Anthropic did not finish Mythos Preview and then begin asking whether release was safe. The company tested the system rigorously, concluded that the surrounding controls were not yet adequate, and chose to keep the model out of public hands. The governance question came before the deployment question, not after.



In most enterprise environments, the order is reversed. Speed is rewarded, and the supporting governance ecosystem has not yet caught up to the pace of capability. Writing on the day of the Glasswing announcement, New York Times columnist Thomas Friedman compared what Mythos Preview represents to the emergence of nuclear weapons and the international nonproliferation efforts that followed, framing it as a capability no single company or country can manage on its own. He has a point, but the civilizational frame should not let any individual organization off the hook. Every team deploying AI today is facing a smaller version of the same question Anthropic answered: are the guardrails in place strong enough for what this system can actually do?





Many leaders cannot yet answer that with confidence. The frameworks, technical standards, and regulatory guidance that would let them answer it rigorously are still being assembled.

This is not a problem to schedule for later

Project Glasswing is a meaningful start. It brings together dozens of organizations, applies a defensive mandate, and is backed by a \$100 million commitment focused on a specific category of risk. But it is not a comprehensive answer to the broader challenge it has put on display. That challenge belongs to every organization that is building or deploying AI.

01

The practical implications are concrete. Treat the adequacy of your guardrails as a precondition for deployment, not a fix-it-later item once something goes wrong. Continuously measure the gap between what your governance documents promise and what your AI systems actually do in production. And remember that the controls designed for today's capabilities will need to be re-examined as those capabilities advance, because static governance will fall behind a moving target.

02

What Anthropic demonstrated with Mythos was rare in any industry: the discipline to ask the hard governance question honestly and to act on the answer, even when the answer was inconvenient. The organizations that will end up on the right side of AI's history are the ones asking that same question now, before some headline-grade incident forces it on them.

