

False Confidence: Why Most Enterprises Are Less Prepared for AI Risk Than They Believe



Across industries, enterprise leaders are expressing growing confidence in their ability to govern artificial intelligence. Surveys show that a majority of decision-makers believe their organizations have adequate controls in place, that they would detect a misbehaving AI model quickly, and that their governance frameworks are keeping pace with deployment. But a closer look at how these same organizations are actually structured reveals a very different picture. What many enterprises have built is not a governance strategy. It is an illusion of one.

01

Recent research surveying enterprise companies with 100 or more employees found that 72 percent of organizations claim to have two or more AI platforms that they each consider their “primary” layer. These platforms span the full range of major providers, from hyperscale cloud and AI lab offerings like Microsoft Azure, Google, OpenAI, and Anthropic to large application vendors like Workday, ServiceNow, and Epic. The result is not a coherent AI strategy. It is a collection of overlapping, often contradictory systems, each expanding the organization's attack surface at a time when AI-driven threats are growing more sophisticated by the month.

The Platform Sprawl Problem

The proliferation of AI platforms inside large enterprises did not happen by accident. It is the predictable outcome of two forces colliding: software vendors racing to embed AI into every product they sell, and enterprise customers rushing to adopt AI before their competitors do. The result is a landscape where individual departments, teams, and business units are each running their own AI initiatives on different platforms, often with minimal coordination.

01

Consider the challenge facing a major hospital system with tens of thousands of employees. After an initial period of experimentation, the organization found itself managing an uncontrolled number of internal AI proof-of-concept projects that had sprouted up across departments. Employees, excited by the possibilities of AI, had launched projects without centralized oversight. The organization eventually had to shut many of them down and take a step back. Rather than continuing to build custom solutions in-house, leadership decided to wait for its existing software vendors to deliver on their own AI roadmaps. The logic was sound: these vendors have enormous resources and have made AI their top priority, so it made little sense to duplicate their efforts internally.

For more content like and follow me:



@bertblevins

INCGPT.COM

But this approach created its own problems. Each of those vendors is now building AI agents that operate differently, follow different protocols, and integrate with the organization's systems in different ways. The enterprise is left needing to build a control plane that can coordinate and orchestrate all of these agents, a layer of infrastructure that did not exist before and that few organizations have figured out how to implement effectively.

The Trap of Easy Beginnings

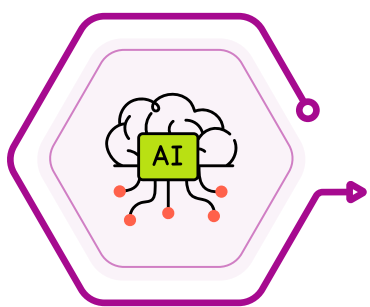
One of the most deceptive aspects of enterprise AI adoption is how simple it appears at the outset. Getting started requires little more than a credit card and an API key. A small team can spin up a working prototype in days, demonstrate impressive results to leadership, and generate enthusiasm throughout the organization. But this initial ease masks the complexity that follows.



The real costs emerge when organizations try to move from pilot projects to production systems. Suddenly, questions of security, compliance, data governance, model monitoring, and cross-platform integration become urgent. The technical debt accumulated during the rapid prototyping phase becomes visible, and organizations discover that the tools they built inside one cloud provider's ecosystem are difficult or impossible to move to another. They have, in effect, built their entire AI capability inside a proprietary cage without realizing it until they try to leave.

Overconfidence in Detection and Control

Perhaps the most concerning finding from recent enterprise surveys is the gap between perceived and actual governance capabilities. More than half of enterprise decision-makers surveyed said they were “very confident” they would detect a misbehaving AI model. Yet nearly a third of those same organizations have no systematic mechanism to identify AI misbehavior until it surfaces through user complaints or external audits. In an environment where telemetry leakage accounts for more than a third of generative AI incidents and the average cost of a data breach has reached \$4.4 million, discovering problems after the damage has already occurred is not governance. It is crisis management.



This overconfidence is compounded by the speed at which AI systems can cause harm. Unlike traditional software, which typically fails in predictable ways, AI models can drift, hallucinate, or behave unpredictably under conditions that are difficult to anticipate. Without continuous monitoring, real-time evaluation, and the ability to intervene instantly, organizations are flying blind while telling themselves they have full visibility.

The Missing Kill Switch

In high-stakes industries like healthcare, the absence of robust control mechanisms is not just a governance gap. It is a safety issue. Industry leaders have been increasingly vocal about the need for what amounts to an emergency stop capability for AI systems, a mechanism that allows an organization to immediately shut down an AI model or agent the moment it begins behaving in ways that could cause harm. Without this kind of hard-stop capability, deploying AI in operational settings is inherently reckless.



This is not a fringe position. The security community group OWASP has formally called for kill-switch capabilities as part of a recommended security framework for AI systems. Yet many organizations continue to deploy AI agents without any such mechanism in place, relying instead on after-the-fact reviews and manual intervention processes that cannot keep pace with the speed at which AI operates.

Why Multiple Primary Platforms Signal a Deeper Problem

When an organization identifies two or more AI platforms as “primary,” it is not demonstrating a sophisticated multi-vendor strategy. It is revealing that it has not made the fundamental decision about who owns the control and security plane. Each additional platform introduces its own authentication model, its own data handling practices, its own agent behavior, and its own set of vulnerabilities. Without a unified governance layer sitting above all of these platforms, there is no single source of truth for what the organization’s AI systems are doing, what data they are accessing, or how they are making decisions.



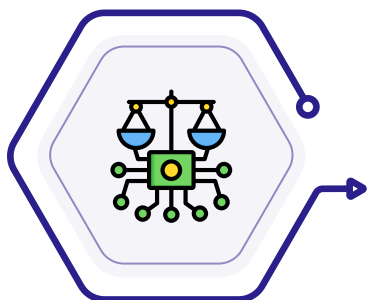
The enterprises that will navigate this landscape successfully are the ones that recognize this conflict of interest for what it is and make deliberate architectural choices about where control lives. That means building or adopting an orchestration and governance layer that is independent of any single vendor, one that provides visibility, policy enforcement, and intervention capabilities across every AI platform the organization uses.

From Governance Theater to Real Accountability

Closing the gap between perceived governance and actual governance requires more than purchasing another tool or adding another policy document. It requires a fundamental shift in how organizations think about AI accountability. Governance cannot be a checkbox exercise conducted at the beginning of a project and then forgotten. It must be embedded into the entire lifecycle of every AI system, from initial design through deployment, operation, and eventual retirement.



This means investing in real-time observability that tracks model behavior continuously, not just at deployment. It means establishing clear ownership of the control plane so that when something goes wrong, there is no ambiguity about who is responsible for shutting it down. It means treating AI agents as identity-bearing entities within the organization’s security architecture, with their own permissions, audit trails, and access controls. And it means accepting that the speed and enthusiasm that characterized the early phase of enterprise AI adoption must now be tempered by the discipline required to operate these systems safely at scale.



The enterprises that claim confidence in their AI governance today may be the ones most at risk tomorrow. The gap between what organizations believe about their preparedness and what the evidence actually shows is not a minor discrepancy. It is a structural vulnerability, and it will remain one until enterprises stop treating governance as an afterthought and start treating it as the foundation on which everything else depends.

