

AI Usage Belongs in Workflow Governance, Not Blank-Check Access

Why giving people unlimited, unmonitored AI is not adoption maturity — and how to govern AI where the work actually happens

Marco Policani

Enterprise Portfolio · PMO · AI Operating Governance

policani.net · linkedin.com/in/marcpolicani

July 2026

DECISION SIGNAL

Govern AI where work happens: by workflow, reliance boundary, and accountable owner.

Executive premise

Many organizations have concluded that the way to accelerate AI adoption is to remove friction: give everyone access to capable tools, do not get in the way, and let a thousand uses bloom. It feels like maturity, the confident opposite of the cautious organization still piloting in a corner. It is not maturity. Blank-check access, unlimited, unmonitored, unscoped AI use, is not governance at all. It is the absence of governance rebranded as enablement, and it produces exactly the ungoverned usage that leads to the incidents, the wasted spend, and the loss of trust that stall AI programs.

The argument of this paper is that AI usage belongs inside workflow governance: it should be governed where the work actually happens, in terms of which workflow the AI is used in, for what work, with what controls, what monitoring, and what limits. This is the opposite of both blank-check access and blanket prohibition. It is the ordinary discipline any organization applies to a capable production capability, applied to AI usage rather than avoided in the name of not slowing anyone down.

The stakes rose with agentic AI. As McKinsey observes for the shift to the agentic era, the risk expands from an AI saying the wrong thing to an AI taking the wrong action, misusing tools, or operating beyond its guardrails, which is precisely why usage cannot be a blank check.¹ Ungoverned usage of a system that only talks is a manageable problem; ungoverned usage of a system that acts is an operational exposure. The rest of this paper defines what workflow governance of AI usage means, why blank-check access fails, and the controls that make usage governable without making it impossible.

Why blank-check access feels like progress but is not

Blank-check access is seductive because it produces immediate, visible activity. Usage numbers climb, people report time saved, and the organization can tell a story of rapid adoption. The problem is that activity is not the same as governed value, and unmonitored activity is not even visible as value. An organization that has given everyone unlimited AI and is not watching how it is used knows only that the tools are being used, not whether they are being used well, safely, or for work that matters.

The deeper failure is that access is not the unit that needs governing. Granting a license answers who may use AI; it says nothing about in which workflow, for what decision, with what oversight, and within what limits, which are the questions that actually determine whether a given AI use is safe and valuable. Governing access while ignoring usage is like issuing everyone a company vehicle and declaring the fleet governed because you tracked who has keys. The keys are not the risk. What happens on the road is the risk, and blank-check access looks away from exactly that.

Blank-check access also tends to degrade into two bad outcomes at once. Where the work is low-stakes, it produces sprawl, a proliferation of unmonitored uses that are individually harmless and collectively impossible to see or improve. Where the work is high-stakes, it produces exposure, because the same blank check that let someone summarize a memo also let someone route a decision or trigger an action with no oversight designed for it. An organization cannot tell these apart without governing usage, which is the thing blank-check access specifically declines to do.

What workflow governance of AI usage means

The alternative can be defined precisely:

Workflow governance of AI usage means governing AI at the point of use, so that for each workflow where AI is applied, the organization has decided what the AI is used for, what controls and oversight apply, what usage is monitored, and what limits are enforced, with the intensity of governance proportioned to the stakes of the workflow. It is neither blank-check access nor blanket prohibition; it is contextual governance that follows the work.

The key word is contextual. The same AI capability used to draft an internal note and used to take an action with customer consequences are not the same governance object, even though they are the same tool, because the workflows differ in stakes. Governing usage means the low-stakes use stays light and the high-stakes use carries the oversight its consequences demand, and the organization can tell which is which. This is exactly the logic the NIST AI Risk Management Framework builds in: its govern, map, measure, and manage cycle frames AI governance as contextual and use-based, driven by how and where a system is used rather than by a single blanket policy.²

From access control to usage governance

The shift this paper argues for is from controlling access to governing usage, and the distinction is operational, not semantic. Access control is a gate at the front: who is allowed in. Usage governance is a discipline throughout: how the capability is used once inside, in which workflows, under what oversight. Organizations are generally good at the first and largely absent at the second, which is why so many have confident access policies and no idea how their AI is actually being used.

Moving to usage governance does not mean surveilling every keystroke or approving every prompt, which would recreate the friction blank-check access was reacting against. It means knowing, at the level of workflows rather than individuals, where AI is being used for consequential work, ensuring those uses have appropriate oversight, monitoring usage so problems and value are both visible, and enforcing limits where unbounded use is a risk. The unit of governance is the workflow-and-its-stakes, not the person and their license.

The controls that make usage governable

Workflow governance of AI usage rests on a small set of controls, each ordinary for a production capability and each typically missing under blank-check access.

- Scoped access to the work, not blanket access to the tool: AI is enabled for the workflows where it has been designed to help, with higher-stakes workflows carrying more deliberate scoping.
- Monitoring and observability: the organization can see how AI is being used across workflows, so both emerging problems and realized value are visible rather than invisible.
- Usage limits as guardrails: bounds on volume, rate, or autonomy that keep an unbounded or misconfigured use from causing outsized harm — an operational safety control, not a budgeting mechanism.
- Evidence and logging: consequential AI uses leave a trail sufficient to reconstruct what was done, so usage can be audited and failures diagnosed.
- Escalation and human oversight proportioned to stakes: high-consequence workflows route uncertain cases to humans and keep a named owner accountable for the AI's use in that workflow.
- Review of usage patterns: someone looks at how AI is actually being used, to catch drift, retire uses that are not working, and extend the ones that are.

None of these controls prohibits AI use; each makes it governable. Together they are the difference between an organization that can answer how is our AI being used, and is it safe and valuable, and one that can only answer how many licenses we bought. The controls are proportioned: a low-stakes workflow may need little more than light monitoring, while a high-stakes one needs the full set. Proportioning is the discipline; applying everything everywhere is as much a failure as applying nothing.

Monitoring is the minimum

If an organization adopts only one element of usage governance, it should be monitoring, because monitoring is what converts blank-check access from a black box into a governable system. The NIST Generative AI Profile treats continuous monitoring, structured feedback, and incident response as core

operating actions across the AI lifecycle, and attention to the human-AI configuration itself as a first-class concern.³ Monitoring is the mechanism that makes all of that possible: without it, an organization cannot detect misuse, cannot see value, cannot diagnose failures, and cannot improve, because it cannot see what is happening.

Monitoring usage is not the same as distrusting users. It is the same instrumentation any serious production system carries, applied to AI. An unmonitored production database would be considered negligently operated; an unmonitored AI capability that people are acting on should be regarded the same way. The organizations that govern AI usage well are not the ones that watch their people; they are the ones that instrument their systems, so that usage is visible at the level of workflows and patterns, and problems announce themselves rather than surfacing as incidents.

Reading the difference

Whether an organization is governing AI usage or merely granting access shows up in what it can answer. The table pairs the blank-check signal with the workflow-governance reality.

Blank-check signal	What it actually tells you	Workflow-governance reality
Everyone has access	Who holds a license	Which workflows use AI, for what, with what oversight
Usage numbers are high	The tools are being used	Whether the uses are safe, valuable, and monitored
No one is complaining	Nothing has gone loudly wrong yet	Problems are visible early through monitoring
We move fast on AI	Friction was removed	Governance is proportioned to stakes, not absent
We trust our people	Oversight was skipped	Consequential uses have oversight and a named owner

The right column is what governed usage lets a leader say. The left column is what blank-check access lets a leader say, and the gap between them is precisely the governance that was skipped in the name of moving fast.

This is not about restricting AI

It is important to be clear about what this argument is not, because the reflex is to hear any call for governance as a call to slow AI down. Workflow governance of AI usage is not restriction; it is the condition for scaling AI safely, and it is frequently what unblocks broader use rather than what constrains it. An organization that can see and govern how AI is used can confidently extend it into higher-stakes workflows, because it has the monitoring and controls to do so responsibly. An organization on blank-check access cannot, because it has no way to know whether the extension is safe, so it either stalls at the safe periphery or lurches into exposure.

Governed usage is therefore an enabler, not a brake. It replaces the false choice between reckless blank-check access and paralytic prohibition with a third option: contextual governance that lets AI go where it adds value, with oversight proportioned to the stakes. That is the posture that scales, and it is available only to organizations that decided to govern usage rather than to look away from it. Removing friction is not the same as enabling AI; enabling AI is giving it a governed path into the work that matters.

Evidence gaps this paper cannot close

Two limits are worth naming. First, this paper argues a governance posture rather than reporting a measured outcome; the claim that governed usage scales better than blank-check access is grounded in the operating logic and in the direction of the trust and risk literature, not in a controlled study comparing the two approaches, which does not yet exist in a transferable form. Second, the right intensity of usage governance is genuinely contextual and cannot be specified as a universal setting: over-governing low-stakes use recreates the friction blank-check access was reacting against, and under-governing high-stakes use is the exposure this paper warns of. The discipline is proportioning, and proportioning requires judgment about stakes that this paper can frame but not pre-compute for any particular organization.

The shadow-AI problem

Blank-check access has a quieter cousin that makes usage governance even more urgent: shadow AI, the use of AI tools the organization did not sanction and cannot see. When official access is either prohibited or unbounded, people route around the gap. Prohibition drives usage to personal tools outside any control, and blank-check access blurs the line so thoroughly that no one can distinguish sanctioned from unsanctioned use. Either way, the organization ends up with consequential work flowing through AI it does not govern and often does not know about.

Workflow governance is the practical answer to shadow AI, because it competes with the shadow rather than merely forbidding it. When the sanctioned path is genuinely useful, scoped to the work, and not burdened with pointless friction, people use it, and the organization gains visibility as a byproduct of providing a better option. Prohibition creates shadow AI by leaving a need unmet; blank-check access hides shadow AI by making all usage look the same; governed usage reduces it by offering a sanctioned path that is both safe and worth using. The goal is not to catch people but to make the governed path the obvious one.

This reframes a security instinct into an operating one. The temptation is to treat shadow AI as a policing problem, to be solved with detection and prohibition. The more durable solution is to treat it as a product problem: the sanctioned, governed path lost to the shadow because the shadow was more useful, and the fix is to make the governed path the better tool, not merely the permitted one. Organizations that internalize this get visibility and safety together, because people choose the path that is governed because it is also good.

Usage limits are guardrails, not budgets

One control deserves special clarification because it is easily misread: usage limits. In a workflow-governance frame, limits on volume, rate, or autonomy are operational safety guardrails, not cost-control measures, and conflating the two undermines both. A rate limit that keeps a misconfigured automation from taking ten thousand actions before anyone notices is a safety control; it happens to also bound cost, but its purpose is to contain harm. Treating that same limit as merely a budget line invites the wrong debate, whether the organization can afford more, when the real question is whether an unbounded use is safe.

Framed as guardrails, limits become part of the same design conversation as escalation and oversight: how much can this AI do, in this workflow, before a human is in the loop or a ceiling is hit? A high-autonomy use in a high-stakes workflow warrants tight bounds and fast escalation; a low-stakes use can run with generous limits and light monitoring. The point is that the bound exists and is set deliberately by the stakes of the workflow, so that no single misconfiguration, prompt injection, or runaway loop can cause damage disproportionate to the use. Blank-check access, by definition, sets no such bound, which is why a single bad configuration under blank-check access can do unbounded harm before anyone sees it.

Keeping the framing operational also keeps this paper honest about its lane. The governance argued for here is technology and operating governance, not financial management; limits, monitoring, and controls are about safe, valuable operation of a capability, and any cost benefit is a side effect, not the point. An organization that governs AI usage only to control spend has missed the larger risk, that ungoverned usage of a system that acts is an operational exposure long before it is a budget problem.

Governing agentic usage specifically

Everything in this paper intensifies when the AI does not merely produce output but takes action, which is the direction agentic systems are moving. McKinsey's framing is explicit that the agentic shift expands risk from wrong statements to wrong actions, tool misuse, and operation beyond guardrails.¹ Blank-check access to a system that only drafts text is a bounded problem; blank-check access to a system that can act, transact, or trigger downstream processes is an open-ended one, because the space of things that can go wrong is the space of actions the agent can take.

This means the workflow-governance discipline is not optional for agentic AI; it is the precondition for deploying it at all. An agent operating in a workflow needs scoped authority to that workflow, monitoring of what it actually does, limits on its autonomy, an evidence trail, escalation for the cases it should not decide, and a named human owner, exactly the usage-governance controls this paper describes, applied with the seriousness that action rather than speech demands. Organizations that learned to govern AI usage while it was mostly generating text will be ready when it starts acting. Those that ran on blank-check access will discover the gap the hard way, at the moment an ungoverned agent does something no one authorized and no one was watching for.

What it looks like in practice

The safest illustrations are operating patterns, not named systems. In governance work, the durable pattern is that AI usage is treated the way any capable production capability is treated: enabled where it helps, instrumented so it is visible, and bounded where it is consequential, with the intensity following the stakes of the workflow rather than a blanket policy.

In one pattern, an organization mapped where AI was being applied and sorted those uses by stakes rather than by tool. Low-stakes drafting uses were left light, with monitoring but little friction; high-stakes uses that fed decisions or triggered actions were given scoped access, oversight, limits, and a named owner. The result was faster adoption where it was safe and real control where it mattered, which blank-check access could not have produced because it could not tell the two apart. Governing by stakes, not by tool, was the move.

In a second pattern, an organization discovered a proliferation of unmonitored AI uses only after an output that no one was watching caused a downstream problem. The corrective was not to prohibit AI but to instrument it: to make usage visible at the workflow level, add oversight to the consequential uses, and set guardrails where autonomy had been unbounded. Once usage was governed, the same AI that had been a hidden exposure became a monitored, extensible capability. The lesson repeats: the risk was never that people used AI; it was that no one could see or bound how.

Building the usage-governance layer

For a leader who accepts the argument, the practical question is how to build usage governance without recreating the friction that made blank-check access attractive in the first place. The answer is to build it as a thin layer organized around workflows and stakes, not as a heavy approval process organized around individuals and tools. The starting move is visibility: instrument AI usage so the organization can see, at the level of workflows and patterns, where AI is being applied and for what. Visibility is the foundation because every other control depends on it; an organization that cannot see its usage cannot proportion its governance.

With visibility in place, the next move is to sort uses by stakes and apply governance accordingly. The large majority of uses are low-stakes and need little beyond monitoring, which keeps the layer light where lightness is safe. The small number of high-stakes uses, those that feed consequential decisions or take real-world actions, get the fuller set of controls: scoped access, oversight, limits, evidence, escalation, and a named owner. This concentration of effort is what makes the layer affordable and non-intrusive; the mistake is to govern everything equally, which is as wasteful on the low end as it is insufficient on the high end.

The final move is to make the governance continuous rather than a one-time setup, because usage evolves. New workflows adopt AI, existing uses change stakes as they scale, and yesterday's low-stakes drafting tool becomes today's input to a decision. A usage-governance layer that reviews patterns on a cadence catches these shifts, moves oversight to where the stakes moved, and retires controls that are no longer needed. This is the same ongoing, contextual posture the NIST framework describes, applied to usage: govern, monitor, and adjust as the use changes, rather than classifying everything once and assuming it stays put.²

Built this way, the usage-governance layer is not a bureaucracy bolted onto AI adoption; it is the operating system that makes adoption safe to accelerate. It gives leaders the answer blank-check access cannot: not just that AI is being used, but that its use is visible, proportioned to stakes, and owned, which is the condition under which an organization can confidently push AI into more of the work that matters.

The shift

The mature posture toward AI is not the confident removal of all friction, and it is not the cautious refusal to let anyone use it. Both are failures of governance, one by omission and one by paralysis. The mature posture is to govern AI where the work happens, proportioning oversight to the stakes of the workflow, monitoring usage so it is visible, and enforcing the limits that keep an unbounded use from becoming an unbounded risk.

Blank-check access mistakes the absence of governance for the presence of trust, and it pays for the mistake later, in incidents no one saw coming and value no one could prove. Govern AI usage as what it is, a capable production capability operating inside real workflows, and the organization gets both the speed it wanted and the control it needs. Stop handing out blank checks. Start governing the work.

Notes and sources

1. McKinsey & Company, "The state of AI trust in 2026: Shifting to the agentic era," 2026. Verified July 7, 2026. <https://www.mckinsey.com/capabilities/tech-and-ai/our-insights/tech-forward/state-of-ai-trust-in-2026-shifting-to-the-agentic-era>
2. National Institute of Standards and Technology, "AI Risk Management Framework Core," excerpt from AI RMF 1.0, 2023. Verified July 5, 2026. <https://airc.nist.gov/airmf-resources/airmf/5-sec-core/>
3. National Institute of Standards and Technology, "Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile," NIST AI 600-1, July 2024. Verified July 9, 2026. <https://doi.org/10.6028/NIST.AI.600-1>

About the author

Marco Policani is an enterprise portfolio, PMO, and AI operating-governance leader. He builds portfolio governance systems where AI carries the analytical load and named humans own every decision — an approach documented across case studies, walkthroughs, and working governance guides on his portfolio site.

Portfolio & governance library: policani.net/governance · Full portfolio: policani.net · LinkedIn: linkedin.com/in/marcpolicani

This white paper may be shared freely with attribution. © 2026 Marco Policani.