

The Real Agent Bottleneck Is Permissions, **Not Intelligence**



Enterprise AI agents are stalling well short of their promise. The cause isn't a weak model — it's the unanswered question of what each agent is allowed to do, on whose authority, and how the system can prove it.

For the past two years, the conversation about enterprise AI has fixated on model capability. Which model reasons best, which one is cheapest per token, which one tops the latest benchmark. Yet the companies actually trying to put autonomous agents into production keep running into a wall that has nothing to do with how smart the model is. The wall is permissions.

An agent that can plan a multi-step task flawlessly is useless — or worse, dangerous — if no one can say with confidence which records it may read, which actions it may take, and whose authority it is acting under. Once an organization moves past demos and into real workflows, the same question surfaces every time: what is this agent permitted to touch, and how does the system know?

Why Capability Was Never the Hard Part

Modern models are already good enough to drive meaningful work. They can interpret a request, break it into steps, call tools, and adjust when something goes wrong. The constraint isn't intelligence; it's trust that can be enforced. An agent acting on a person's behalf inherits a thorny problem that human employees solve through years of accumulated context, role definitions, and judgment — none of which an agent has by default.



The result is a familiar failure pattern. Teams wire an agent directly to raw data to get it moving quickly, and in doing so they strip away the carefully built security model that governed that data in the first place. The agent technically works, but its access is far too broad, its actions are hard to audit, and its results can't be trusted for anything consequential.

For more content like and follow me:

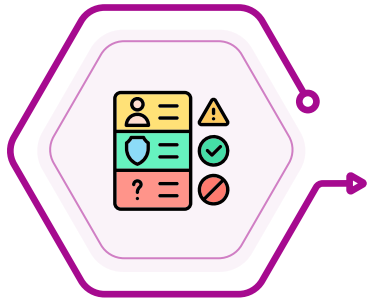


@bertblevins

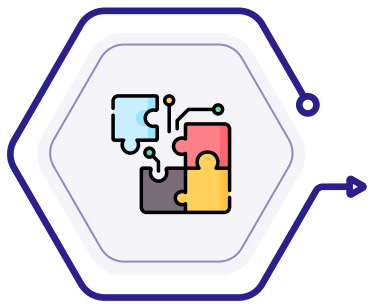
INCGPT.COM

Permissions Belong Where the Data Lives

A growing consensus among practitioners is that permissioning cannot be bolted on from the outside. If the rules about who can access what are defined somewhere other than the system that actually holds the data, the gap between the two becomes the vulnerability. The moment an agent reaches around that boundary, the integrity of approvals and access controls is already compromised.



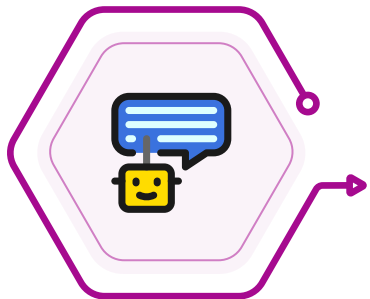
This is why several vendors are converging on the same architectural answer: make the existing system of record the governance layer for agents, rather than inventing a parallel permissions scheme that has to be kept in sync. When an agent's authority is grounded in the same place the data is stored and the same approval logic humans already use, the security model stays intact instead of being quietly discarded.



Workday has pushed this idea explicitly with Sana, its agent system of record introduced earlier in the year, which is meant to ensure that the integrity of an organization's approval and security model is always preserved when agents act. The company has also extended its work with Google so that agents grounded in Sana can operate within Gemini Enterprise — a signal that the system-of-record-as-governance approach is being designed to span platforms rather than live in a single silo.

Inheriting Access, Not Inventing It

The cleanest implementations treat the agent as an interface layer rather than a privileged actor. The agent doesn't get its own sweeping credentials; it operates with the identity and token of the person using it, so it can reach exactly what that person could reach and nothing more. Access is inherited, not granted anew — which means the existing access controls keep doing their job instead of being bypassed.



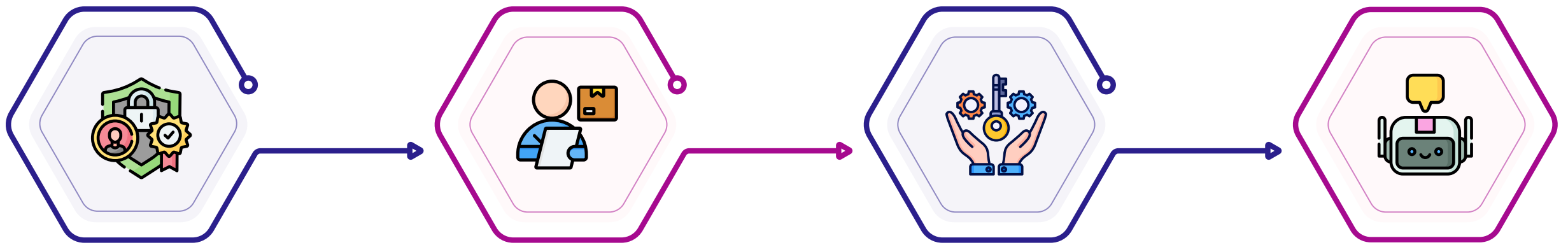
That philosophy pairs naturally with what one large AI operator described as “dumb guardrails” — strong, simple, hard access controls rather than clever alignment tricks layered on top of over-broad permissions. The unglamorous controls are the ones that hold up in production.

Ownership Is the Other Half of the Problem

Permissions answer what an agent may do. Ownership answers who is accountable when it does it. Without clear ownership of an agent's performance, its costs, and its actions, autonomous systems drift toward chaos: nobody is sure which agent took which action, why it cost what it cost, or who should intervene when it misbehaves.

Mature deployments therefore pair access control with accountability:





Bounded authority:

each agent's permissions are explicit, scoped to a real identity, and enforced by the system that owns the data.

Auditability:

every action an agent takes is traceable to a person, a permission, and a moment in time.

Clear ownership:

someone is responsible for each agent's behavior, spend, and outcomes — not the abstract “the AI did it.”

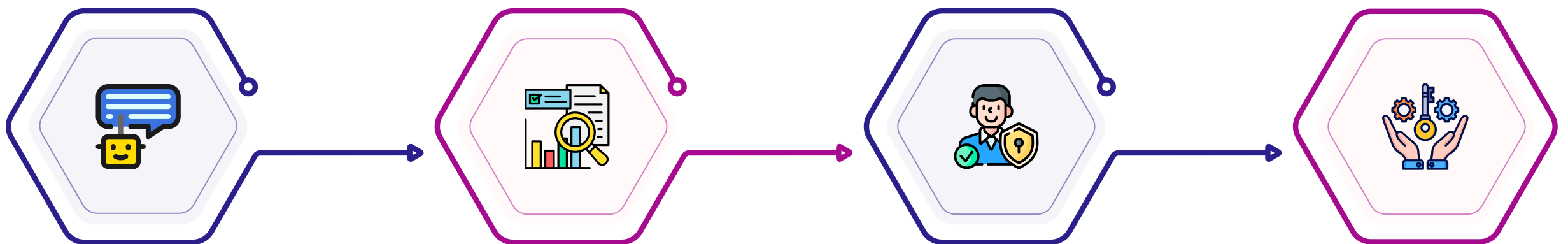
Fail-safe defaults:

when context is missing or authority is unclear, the agent does less, not more.

What This Means for Enterprise Buyers

The practical takeaway is that the next phase of enterprise AI will be won on governance, not on leaderboard scores. Choosing an agent platform is increasingly a question of how it handles identity, permissions, approvals, and audit — because those are the things that decide whether an agent can be trusted with anything that matters.

Teams evaluating agentic systems should press on a few concrete points:



Does the agent inherit a user's real permissions, or does it get its own broad access?

Where are permissions defined — inside the system that holds the data, or in a separate layer that can fall out of sync?

Can every agent action be traced back to an authorized person and an explicit approval?

Is there clear ownership of each agent's actions, costs, and failures?

The Bottom Line

The industry spent its energy making agents smart. The work that remains — and the work that will actually unlock production use — is making them governable. As long as permissioning is an afterthought, even the most capable agent is stuck at the demo stage. Solve the question of what an agent is allowed to do and on whose behalf, and the model was never really the bottleneck at all.

