

Earning the Handoff:

Where Enterprises Trust AI Agents and Where They Hold Back



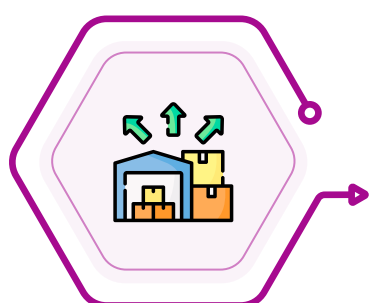
Corporate spending on artificial intelligence is climbing fast, and 2026 is shaping up to be the year boards stop asking whether AI works and start asking what it returns. Gartner has labeled this an inflection point, the moment organizations are expected to tie their AI initiatives directly to business strategy rather than treating them as experiments. With analysts forecasting AI budgets to swell by nearly half this year alone, the pressure on executives to show real financial results has never been sharper. That pressure is steering attention toward a particular flavor of the technology: agentic AI, the kind that does not just answer questions but takes action.



Nowhere is that opportunity more obvious than inside the technology function itself. Infrastructure costs are on a trajectory to double or triple by the end of the decade, even as the budgets meant to cover them stay flat. Faced with that squeeze, the people who build and maintain enterprise systems, engineers, developers, architects, and platform specialists, have spent the past year and a half quietly putting agents to work. The question is no longer whether they will adopt the technology, but how far they are willing to let it run on its own.

Confidence, Ranked

A recent study set out to measure exactly that. Researchers surveyed three hundred technology experts around the world and asked them to rate their confidence in handing one hundred and one distinct tasks over to AI agents across AI, data, and cloud workflows. The resulting ranking offers something more useful than a general optimism reading. It shows, task by task, where practitioners are comfortable stepping back and where they insist on keeping a hand on the wheel.



The headline takeaway is that confidence is high and rising, but it is not evenly distributed. Agents earn the most trust when the work is measurable and repeatable. Drafting reports, producing boilerplate code, and clearing away routine, repetitive chores are the kinds of jobs technologists are happy to delegate, because the output is easy to check and the cost of a mistake is low. As the work shifts toward multistep reasoning and genuine judgment, confidence does not vanish, but it cools. That is the frontier the report describes: the line between tasks an agent can finish reliably and tasks that still demand a human to weigh the call.

For more content like and follow me:

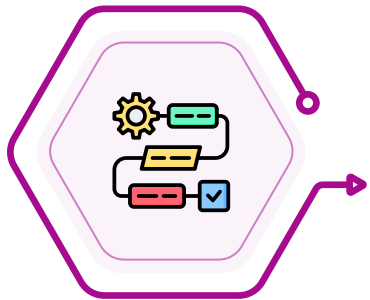


@bertblevins

INCGPT.COM

Why Data Workflows Lead the Pack

The single most trusted domain turned out to be data. Tech teams place the most faith in agents where structure already exists to anchor a decision, and data work is full of that structure. Monitoring data quality, flagging anomalies in dashboards, watching real-time streams for irregularities, and profiling datasets all topped the confidence rankings. The common thread is that the experts closest to where the data is generated can hand an agent enough context to act sensibly and produce a result they can trust.



That detail matters more than it might first appear. The reason data workflows score so well is not that the tasks are simple. It is that the surrounding information is rich and close at hand. When an agent operates inside a domain where the ground truth is well defined, it has the footing it needs to perform without constant supervision.

The Real Bottleneck Is Context

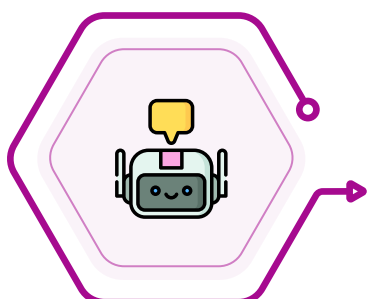
Where readiness falls off, the culprit is almost always missing business context. The more intricate a task becomes, the more reasoning an agent has to do, and the more it depends on understanding the specific circumstances of the organization it serves. Today that context is the hard part. The tooling that feeds enterprise knowledge into an agent at the speed and quality decision-makers need is still immature, especially when the relevant data is scattered, messy, or locked inside systems that were never designed to talk to one another.



This is the gap that separates a promising demo from a production deployment. An agent that can summarize a clean dataset is a long way from an agent that can navigate a company's tangled history of exceptions, legacy decisions, and unwritten rules. Closing that distance is less about raw model intelligence and more about plumbing: connecting the right information to the agent at the right moment, reliably enough to act on.

Trust Is Built on Familiar Guardrails

The experts interviewed for the study expect confidence to keep climbing as teams gain experience and as the surrounding environment matures. A telling observation from the research is that agents start to feel trustworthy when they are designed to operate inside the same boundaries, identity systems, and governance models that human teams already rely on. In other words, an agent earns trust not by being clever but by behaving like the systems an organization has already decided to depend on.

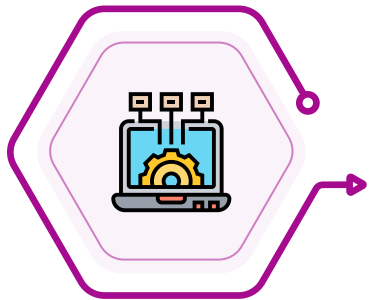


That framing should resonate with anyone responsible for security and identity. An agent is, in effect, a new kind of actor inside the enterprise, one that authenticates, accesses resources, and makes changes. If it operates under the same access controls, the same audit trails, and the same governance discipline applied to human users and service accounts, it slots into an existing model of accountability. If it operates outside those controls, it becomes one more piece of unmanaged automation, the kind that creates risk faster than it creates value.



The Human Stays in the Loop

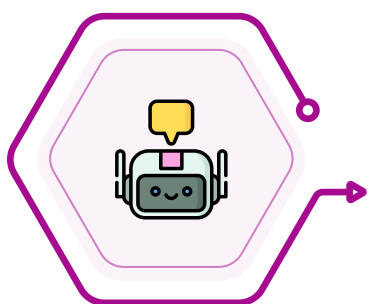
For all the enthusiasm, the report is clear on one point: human oversight remains a deciding factor in whether agentic deployments succeed. Automated decision-making carries real consequences, and teams are right to refuse to hand off work they cannot verify. The goal that emerges from the research is not a workforce of unsupervised agents but a partnership, where agents handle the measurable, high-volume work and people concentrate on the judgment calls, the edge cases, and the strategic direction.



Far from threatening the careers of the technologists who adopt them, agents appear to be reshaping those roles toward higher-value work. The practitioners best positioned to benefit are the ones who learn to design, supervise, and govern these systems, treating the agent as a capable but accountable member of the team rather than a black box to be feared or blindly trusted.

What It Means for Security and Identity Leaders

The path forward suggested by this research is pragmatic. Start where confidence is already high and the stakes are contained, lean on the data-rich workflows where agents perform best, and resist the urge to delegate complex judgment before the context to support it exists. Above all, treat every agent as an identity that must be governed. The organizations that will pull ahead are not the ones that deploy the most agents, but the ones that fold those agents into the trust frameworks they have spent years building, so that scaling automation never means loosening control.



Agentic AI is moving from novelty to infrastructure. The teams that win the next phase will be the ones who understand that confidence is not granted to a technology, it is earned, task by task, through visibility, governance, and a clear-eyed sense of what a machine should and should not be allowed to decide on its own.

