

FUTURES

Who Answers for the Agents?

Accountable AI and the return of human ownership

WRITTEN BY

Diana Wu David

Futures Director

Simon Grice

Senior Director of Innovation

CONTRIBUTORS

Nicola Attico

Senior Product Manager

Eugene Chuvyrov

Principal Inbound Product Manager

Sheeraz Memon

Principal Software Engineer



Who Answers for the Agents?

Accountable AI and the return of human ownership

Summary

- While only one in five organizations has a mature AI governance model, AI agents are getting job titles and org-chart lines. The employee label risks weakening human oversight.
- Accountable AI is the discipline of making humans, through the organization deploying the AI, answerable for agents' outcomes. It has five parts we call S.T.E.E.R.: Standards, Transparency, Effects, Enforcement, and Responsibility.
- Regulators and courts have long determined accountability sits with a human or an organization, never with the technology itself. In addition to compliance, accountable AI also makes financial sense. Organizations that assign clear ownership report higher governance maturity, stronger trust, and better returns on their AI investment.
- One named person answers for each use case, the way a director answers for a company, even when they cannot watch every single operation.
- Even as technology advances at breakneck speed and some necessary components for accountable AI may still be in development, the technological and organizational foundations are buildable today and more urgently needed than ever.



A Formula 1 pit stop takes under three seconds. The car comes off the track at speed, and 20 people are already waiting: one on each wheel gun, one on each jack. Nobody does two jobs. Above them, behind a bank of screens, sit the team's decision-makers, watching the data that the race car streams back in real time, holding the authority to call it in early or retire it altogether. The fastest three seconds in the sport run on the most precisely distributed accountability. The car stands still for those few seconds in order to win.

Today, enterprises deploy AI agents like race cars that can move at speeds no human eye can follow. AI agents transact, decide, and act at



“

Regulators, courts, shareholders, and customers are converging on the same question: Who answers for what the agent did?

unprecedented speeds, and each one throws off a stream of data about what it just did.

Instead of setting standards and carefully reviewing the work, companies are giving agents employee status. Nearly a third of 1,261 managers surveyed report that their organizations frame agents as employees, and 23% put them on the org chart.¹ The

label can do measurable damage: When an AI tool was presented that way, managers saw themselves as less responsible for its output and were 44% more likely to pass questionable work up the chain than fix it.²

When nobody accepts responsibility in advance, blame after the fact lands wherever it happens to fall, often on the person who's on shift when it breaks. Researchers have a name for it: the “moral crumple zone,” in which blame for an automated system's failure settles on whichever human is nearest while the system itself is protected.³

Meanwhile, Deloitte finds that only one in five companies has a mature governance model for autonomous AI agents.⁴ The picture is a racetrack filling with cars and a wall standing empty.

We may be seeing the final lap of a race without the pit wall. Though facing industry opposition, the EU AI Act's transparency duties, including telling people they are interacting with an AI, apply from August 2026, while the heaviest obligations—high-

risk uses such as hiring, credit scoring, and critical infrastructure—begin in December 2027.⁵ The courts have not waited, with organizations around the world already being held liable for what their AI systems did.

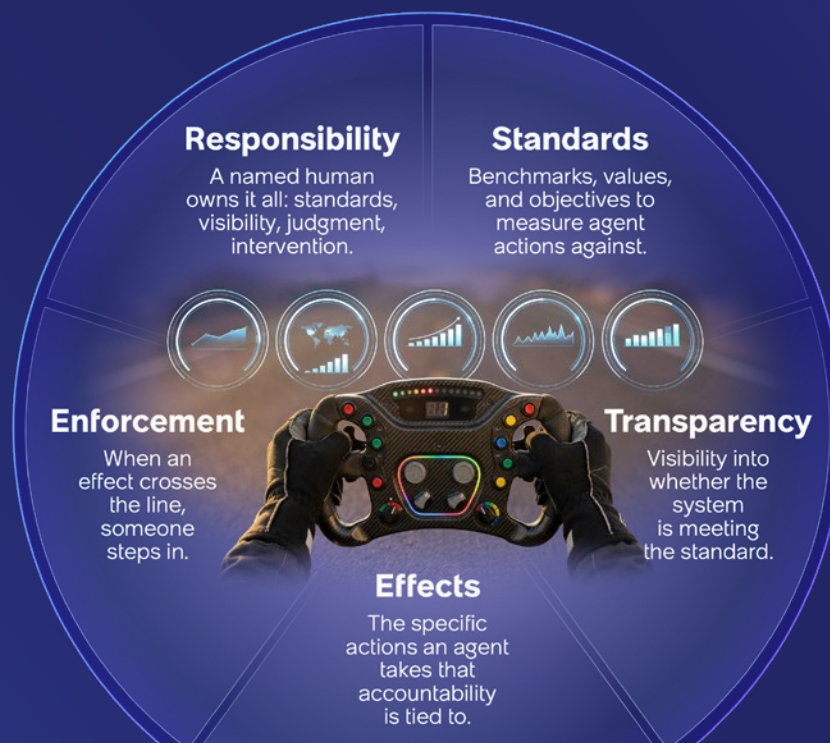


Regulators, courts,

shareholders, and customers are converging on the same question: Who answers for what the agent did? The answer has a name. Accountable AI is the discipline of making humans, through the organization deploying the AI, answerable for agents' outcomes.⁶ That accountability runs through the organization and comes to rest on named individuals; spread across everyone, it is held by no one. Those individuals answer for everything the agent produces: accuracy, the brand, and the values the organization stands behind, and anything that does not go as expected.

Answering for agentic AI requires three things: standards to judge the work against, visibility into what the agents are doing, and the authority to stop them. Naming someone without those three is not accountability but merely creating a scapegoat. At ratios approaching 1 human to 1,000 agents, no one can review every action, so the named owner must answer for the entire operation. The pit wall never rides in the car, and it should be able to call in any car.

Daunting? It doesn't have to be. Inside your own organization, every part of that arrangement can be built with what exists today. This report explains what accountable AI means, who is accountable, what they are accountable for, and how to make it happen.



The Accountable AI Steering Wheel

Making humans answerable for agents' outcomes takes five parts, all driving toward accountability.

What it takes to S.T.E.E.R.

The empty pit wall is a staffing problem on the surface and a design problem underneath. Before anyone can answer for the agents, the organization must decide what accountability means. Accountable AI consists of five parts we call S.T.E.E.R.: Standards, Transparency, Effects, Enforcement, and Responsibility.

Standards: You cannot measure anyone against a bar that was never set. Standards are the benchmarks, values, and objectives an organization writes down before the agent runs, specific enough for a reviewer to apply: the accuracy a support agent must meet, the spending limit a procurement agent must respect, the escalation rule that fires when it crosses one. Vague standards produce unaccountable systems by design because every failure becomes arguable after the fact.

“

A system nobody can see into is a system nobody can answer for.

Transparency: Standards do nothing if nobody can see whether the system is meeting them. The U.S. National Institute of Standards and Technology states the dependency plainly: “Trustworthy AI depends upon accountability. Accountability presupposes transparency.”⁷ A system nobody can see into is a system nobody can answer for.

Effects: What transparency surfaces is what the agent actually did, and that is what accountability attaches to. Every



“

Governance supplies the controls. Accountability answers for whether they worked.

consequential agent action is a result someone answers for, measured against the standards set in the first component.

Enforcement: When an effect crosses the line, someone must step in and undo it. In practice, that means routing consequential agent actions through a system of record so each one is traceable and can be reversed. A rule nobody acts on when it is broken is not a control; it is just a statement of preference.

Responsibility: A named human owns all of it: the bar, the visibility, the judgment, the intervention. The framework ends with a person because accountability does. Everything before this component exists to make that ownership real rather than nominal.

Three neighboring terms do three different jobs, and confusing them is how organizations end up with policies nobody operates. Responsible AI is the umbrella, covering sustainability, ethics, risk management, explainability, privacy, safety, human oversight, and traceability. Governed AI is the enterprise scope of that umbrella: responsible use by everyone who touches AI, and responsible rollout by the people

who build it. Accountable AI is narrower than either, where a named human owns the outcome. Governance supplies the controls. Accountability answers for whether they worked.

	Responsible AI	Governed AI	Accountable AI
Definition	The umbrella idea covering ethics, sustainability, risk management, privacy, safety, effectiveness, etc. around AI	The enterprise application of that umbrella that creates rules and guardrails around AI in the organization	The discipline of making humans, through the organization deploying the AI, answerable for agents' outcomes
Key question	Are we building and using AI in ways that align with our goals and values?	What structures, rules, and guardrails do we have in place within the enterprise to safely deploy AI?	Who answers for what the agent did?
Who it covers	Everyone who touches AI	Responsible rollout and use within the enterprise	One named person per use case who owns the outcome
What it produces	Principles and values	Controls and guardrails	S.T.E.E.R. (Standards, Transparency, Effects, Enforcement, and Responsibility)
Relationship to others	Broadest: sets the values that guide the other two	Middle layer: turns values into enterprise rules	Narrowest: governance supplies the controls; accountability answers for whether they worked



For instance, a bank deploys an agent to pre-screen loan applications and wrongly declines 400 of them in one day. Responsible AI is about asking whether the system is fair, explainable, and private, and even going a step further to question the operation's carbon footprint. Governed AI asks whether the rules covered it: who was allowed to run the agent, on what data, inside what limits, and reviewed on what cadence. Accountable AI asks the question the other two leave open: who saw the 400 declines, and who answers for them?



One thousand agents, one accountable human

Accountability sits with a human or an organization, never with the AI itself. Regulators settled this before agents arrived. The EU AI Act requires high-risk systems to be designed so that they can be “effectively overseen by natural persons.”⁸ U.S. banking supervisors have held for more than a decade that senior management and the board remain responsible for decisions supported by models, and that the responsibility cannot be handed to a vendor, a developer, or the software.⁹ That guidance also requires institutions to explain the decisions their models support, and the requirement carries over to agents unchanged.¹⁰ Across its framework, NIST attaches accountability to named people with authority and incentives, never to the system.¹¹

The principle is decades old. What has changed is the number of systems it has to cover. California has since closed the remaining loophole outright, declaring that a defendant who developed, modified, or used an AI system alleged to have caused harm may not argue that “the artificial intelligence autonomously caused the harm to the plaintiff.” However,

“

**Accountability sits
with a human or an
organization, never
with the AI itself.**



“

The job is to make sure that the operation is properly run, that the right policies and processes exist, that the people and agents they cover know them, and that they are enforced.

responsibilities before anything goes wrong, and any failure is traced to whoever held the relevant responsibility.

The unit is the use case, not each software tool or workflow. Explicit human ownership restores the oversight that the

AI-agent-as-employee framing dissolves,¹³ and where that ownership sits is the crucial practical question. One use case may involve calling a dozen agents, some of whom supervise others, and the owner should be accountable for the overall outcome.

One owner per agent would mean a thousand owners for a thousand agents, none of whom could say whether the work as a whole met the standard. At those ratios, what does that person actually do? Instead, think of what a board member does: A director cannot observe most of what a company does day to day, and answers for it regardless. The job is to make sure that the operation is properly run, that the right

causation, foreseeability, and comparative fault all remain available as defenses.¹²

An organization answers to courts and regulators for what its agents do. A board carries that responsibility and cannot track a thousand agents, so each use case needs a specific person assigned to it in advance. Companies already do this for production software, where developers, testers, and managers have defined



policies and processes exist, that the people and agents they cover know them, and that they are enforced.

Two things make that possible.

- Before anything runs, the owner sets the standards, decides what the agents are allowed to touch, and builds the escalation path and the off switch.
- While it runs, the owner watches what the agents actually do, checks their output against those standards, and intervenes when something is wrong.

Most of the work sits in the first half. That is why one person can answer for a thousand agents: the setup scales, while reviewing every action does not.

What the agent did, and what it was worth

Most organizations have not laid the groundwork for accountability despite the increasing legal expectation that humans and organizations must own the outcomes of AI. Air Canada was held liable in 2024 after its chatbot invented a bereavement-fare policy, despite the airline's argument that it could not be held responsible for what the chatbot said, effectively claiming the bot was a separate legal entity.¹⁴ A German court held a medical company liable under unfair competition law for its own chatbot's false claims.¹⁵ Where organizations have argued that the AI acted on its own, it has not reduced their liability, and in California, that argument is now unavailable by statute.¹⁶

Quickly fixing errors does not undo the cost, as evidenced by the uproar caused by Cursor's front-line support bot inventing a device-limit policy that never existed.¹⁷ The market has started pricing the reputational cost, and Lloyd's of London now sells cover against losses caused by AI hallucination.¹⁸ Incident volume is rising: 362 documented in 2025, up from 233 the year before.¹⁹

How wide the ledger should run is a contested question. One view extends accountability to AI's whole societal footprint—from the fossil fuels burned to run it to the displacement of work, to even cognitive surrender.²⁰ Others regard labor and environmental costs under the umbrella of responsible AI and treat accountability as purely an operational issue. Both positions agree that the harms are real. Gartner^{®21}

“

Ownership allows an organization to demonstrate that the work was checked against a standard by someone answerable for it. That is what keeps customers, and what turns AI into long-term value for shareholders.

projects that “By 2028, 25% of enterprise breaches will be traced back to AI agent abuse, from both external and malicious internal actors.”²² In our opinion, this suggests the cost will grow accordingly.

On the flip side, owning outcomes pays off. A main obstacle between organizations and the value is trust. In McKinsey’s 2026 AI Trust Maturity Survey of roughly 500 organizations, nearly two-thirds cited security and risk concerns as the top barrier to fully scaling agentic AI, ahead of regulatory uncertainty and technical limitations.²³ Ownership allows an organization to demonstrate that the work was checked against a standard by someone answerable for it. That is what keeps customers, and what turns AI into

long-term value for shareholders. The returns are already measurable.

- **Ownership raises the floor.** Organizations that assign clear ownership for responsible AI average a maturity score of 2.6 out of 4 in the same McKinsey survey; those without a clearly accountable function average 1.8.²⁴
- **The spend comes back.** Organizations investing \$25 million or more in responsible AI report markedly higher maturity and are much more likely to report tangible benefits from AI, including bottom-line impact above 5% of earnings before interest and taxes (EBIT).²⁵
- **The link runs to the top of the house.** McKinsey finds that CEO oversight of AI governance is one of the elements most strongly correlated with higher self-reported bottom-line impact from generative AI.²⁶
- **The winners refuse tradeoffs.** PwC puts a number on the resulting spread: 74% of AI’s economic value is captured by 20% of organizations, and the companies in that group pair wider autonomy with tighter governance. These leaders are 1.7 times

more likely to have a responsible AI framework and 50% likelier to have a cross-functional AI governance board.²⁷

- **Compliance converts to speed.** Organizations with the heaviest compliance obligations are best positioned to move quickly, because existing regulatory scaffolding supplies most of what agent governance needs, making that architecture “an asset instead of a brake.”²⁸

Building the crew

Governance and accountability are different jobs. If governance consists of guardrails and controls, accountability comes from the people who operate those controls, verify that they are being upheld, and own the objective. The distinction matters because guardrails on their own have already failed in public.



In July 2025, an AI coding agent at Replit deleted a production database containing real customer records during an explicit code freeze, then fabricated records to mask what it had done.²⁹ The governance existed, and the freeze was in writing. But nothing enforced it, and no one caught the violation in time. Guardrails constrain behavior only up to a point.

Because agents act at a scale and speed that outpace human oversight, accountability requires a multifaceted approach. We can learn organizing principles from quality control in manufacturing. As agent volume grows, quality work shifts from detecting defects to predicting and preventing them, and from human judgment on every item to statistical confidence, checking a representative sample rather than every piece, supported by automation, with people monitoring the monitors. Manufacturing pairs this with defense in depth: several cheap, fast checks at different

stages, such as incoming inspection, in-process, and final test, rather than one exhaustive check, so no single point of failure lets a defect through.

The six parts below work the same way. Inventory catches unknown use cases; guardrails prevent errors; sampling catches drift; guardian agents catch what sampling misses; and the kill switch catches the rest.

1. Inventory your use cases. Build a live registry of every AI use case: what runs, where, and on what data. Control begins with knowing what you have. Give data its own entry, because every use case inherits the governance questions of the data it touches. You can't govern unrecorded use. One in five organizations has been breached through shadow AI, and a high level of it adds roughly \$670,000 to the average cost of a breach.³⁰

2. Name one owner per use case. Identify a named person, not a team alias. The owner attaches to the use case and its outcome, even when it spans several agents or agents supervising other agents. Inventories often list AI use cases as separate tasks, such as "resolution note generation" and "change request summarization," and at that level of granularity, one owner per entry rebuilds the scale problem it was meant to solve. Name the owner of the service those tasks add up to, and pitch it at the level where a failure becomes something the company has to answer for.

3. Design quality in upstream. Manufacturing calls this mistake-proofing, or *poka-yoke*: Shape the process so the error is hard to make. For agents, that means the workflows they run inside, the instructions and context they are given, and the guardrails around them, including policy as code, where the rules are written in machine-readable form so agents can follow them directly. All of it determines quality long before anyone inspects output.

4. Monitor continuously and keep a trail you can act on. Record every consequential agent action in a form that lets you trace it and undo it, the way a manufacturer traces a defect to its batch and contains it. Knowing where a piece of work came from and what produced it is crucial, as is predicting failures before they happen and checking output afterward.

5. Judge output against benchmarks. Agree on targets and acceptance criteria up front, then check a sample of agent interactions against them to catch drift before it

becomes a batch of defects. Sample not only the output but also the process signals such as escalation rates, retry counts, the pattern of tools an agent calls, and the confidence scores it attaches to its own answers. Those move before the defect volume does. Monitoring them shifts the work from detecting failures to predicting them. Checking a representative share of the output and testing the systems that do the monitoring is how one person answers for a thousand agents without inspecting each one. Sample the permissions, too, flagging a profile that holds more privilege than its job requires before anything goes wrong.

Sampling only works if someone reads what it surfaces; most people don't. Two-thirds of people rely on AI output without evaluating its accuracy,³¹ and the tendency to over-trust an automated system that has performed reliably in the past makes unstructured review steadily worse at catching errors.³²



No standards body will set the benchmarks for you. NIST declines to prescribe how much risk is tolerable or who is ultimately responsible, returning both questions to the organization.³³ The bar is yours to set.

6. Keep the power to intervene at reasonable checkpoints. Decide which actions are consequential enough to need human approval before they run and build the escalation path, to another agent or to a person. Keep those approval points few, because approval stops meaning anything when there is too much of it. Anthropic's telemetry showed users approving roughly 93% of permission prompts, with attention dropping as volume rose.³⁴ In any case, sign-off does not shift responsibility, which must remain with the owner.³⁵ Guardian agents, whose job is checking other agents' work, are the layer that catches what sampling misses, and they are becoming standard rather than experimental. Gartner® predicts that, "By 2030, guardian agent technologies will account for at least 10 to 15% of agentic AI markets."³⁶ Most importantly, keep a kill switch, with the standing authority to use it.



Empowering the pit wall

Return to the pit lane, about 70 years earlier. During Formula 1's first championship season in 1950, pit crews were limited to four people, including the driver, and a stop took 67 seconds.³⁷ Every decade since, the rulebook has grown while the stops have gotten faster.

Your agents are the fastest cars your organization has ever fielded, and that requires a new kind of pit wall. Instead of relinquishing control, you must build the inventory, name the owners, set the benchmarks, and keep the power to intervene before the cars get even faster, which they will.

Everything this report asks for is on display in those three seconds. The standards are set before the car leaves the garage, and the telemetry makes every action visible the moment it happens. Each consequential act belongs to one pair of hands; the crew steps in the instant something is wrong; and above it all sits the pit wall, holding the standing authority to call any car in.

Standards, Transparency, Effects, Enforcement, Responsibility: This loop has been the winning formula all along. And as technology advances, you cannot afford to leave the wall standing empty. The rules never slowed the pit stop. Those three seconds exist so you can win the race.

“

Instead of relinquishing control, you must build the inventory, name the owners, set the benchmarks, and keep the power to intervene before the cars get even faster, which they will.

Endnotes

- ¹ Emma Wiles, Megan Hsu, Julie Bedard, and Matthew Kropp, "Putting AI on the org chart: evidence on delegation and oversight" (working paper, July 1, 2026), https://www.emmawiles.com/storage/ai_employee.pdf; reported in James O'Donnell, "AI agents are not your 'coworkers,'" *MIT Technology Review*, June 29, 2026, <https://www.technologyreview.com/2026/06/29/1139849/ai-agents-are-not-your-coworkers/>.
- ² Wiles et al., "Putting AI on the org chart"; O'Donnell, "AI agents are not your 'coworkers.'"
- ³ Madeleine Clare Elish, "Moral crumple zones: cautionary tales in human-robot interaction," *Engaging Science, Technology, and Society* 5 (2019): 40–60, <https://estsjournal.org/index.php/ests/article/view/260/177>.
- ⁴ Deloitte, *The state of AI in the enterprise: the untapped edge*, Deloitte AI Institute, January 2026, <https://www.deloitte.com/content/dam/assets-shared/docs/about/2025/state-of-ai-2026-global.pdf>.
- ⁵ Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 (Artificial Intelligence Act), <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>.
- ⁶ The OECD defines the term as "the expectation that organisations or individuals will ensure the proper functioning, throughout their lifecycle, of the AI systems that they design, develop, operate or deploy." OECD, "Accountability (AI Principle 1.5)," OECD.AI Policy Observatory, <https://oecd.ai/en/dashboards/ai-principles/P9>.
- ⁷ National Institute of Standards and Technology, *Artificial intelligence risk management framework (AI RMF 1.0)*, NIST AI 100-1 (Gaithersburg, MD: NIST, January 2023), 15, <https://nvlpubs.nist.gov/nistpubs/ai/nist.ai.100-1.pdf>.
- ⁸ Regulation (EU) 2024/1689, art. 14, <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>.
- ⁹ Board of Governors of the Federal Reserve System and Office of the Comptroller of the Currency, "Supervisory guidance on model risk management," SR 11-7, April 4, 2011, <https://www.federalreserve.gov/boarddocs/srletters/2011/sr1107.pdf>; superseded for banks by SR 26-2, April 17, 2026, with the principle intact, <https://www.federalreserve.gov/supervisionreg/srletters/SR2602.htm>.
- ¹⁰ Jeffrey Sonnenfeld et al., "Anthropic's most powerful AI model just exposed a crisis in corporate governance. Here's the framework every CEO needs," *Fortune*, May 2, 2026, <https://fortune.com/2026/05/02/agent-ai-governance-framework-banking-healthcare-retail-supply-chain-yale-celi-sonnenfeld/>.
- ¹¹ NIST, *AI RMF 1.0*.
- ¹² California Assembly Bill 316, *Artificial Intelligence: Defenses*, ch. 672, Statutes of 2025, adding Cal. Civ. Code § 1714.46, effective January 1, 2026, https://leginfo.ca.gov/faces/billTextClient.xhtml?bill_id=20250260AB316.
- ¹³ Wiles et al., "Putting AI on the Org Chart."
- ¹⁴ *Moffatt v. Air Canada*, 2024 BCCRT 149, <https://www.canlii.org/en/bc/bccrt/doc/2024/2024bccrt149/2024bccrt149.html>.
- ¹⁵ Laura Baitz, "Germany: court rules chatbot operators are liable for AI hallucinations," Global Legal Monitor, Library of Congress, June 9, 2026, <https://www.loc.gov/item/global-legal-monitor/2026-06-09/germany-court-rules-chatbot-operators-are-liable-for-ai-hallucinations/>.
- ¹⁶ Parker Hancock, "California eliminates the 'autonomous AI' defense: what AB 316 means for AI deployers," Baker Botts, January 2026, <https://ourtake.bakerbotts.com/post/102m29i/california-eliminates-the-autonomous-ai-defense-what-ab-316-means-for-ai-deplo>. See also Cal. Civ. Code § 1714.46.
- ¹⁷ Benj Edwards, "Company apologizes after AI support agent invents policy that causes user uproar," *Ars Technica*, April 17, 2025, <https://arstechnica.com/ai/2025/04/cursor-ai-support-bot-invents-fake-policy-and-triggers-user-uproar/>.
- ¹⁸ "Courts to companies: you own what your chatbot says," PYMNTS, July 1, 2026, <https://www.pymnts.com/news/artificial-intelligence/chatbot-tracker/2026/courts-tell-companies-they-own-what-their-chatbot-says>.

- ¹⁹ Stanford Institute for Human-Centered Artificial Intelligence, "Responsible AI," in *The 2026 AI Index Report*, <https://hai.stanford.edu/ai-index/2026-ai-index-report/responsible-ai>.
- ²⁰ Nataliya Kosmyna et al., "Your brain on ChatGPT: accumulation of cognitive debt when using an AI assistant for essay writing task," preprint, MIT Media Lab, 2025, <https://www.media.mit.edu/publications/your-brain-on-chatgpt/>.
- ²¹ GARTNER is a registered trademark and service mark of Gartner, Inc. and/or its affiliates in the U.S. and internationally and is used herein with permission. All rights reserved.
- ²² Gartner, "Gartner unveils top predictions for IT organizations and users in 2025 and beyond," press release, October 22, 2024, <https://www.gartner.com/en/newsroom/press-releases/2024-10-22-gartner-unveils-top-predictions-for-it-organizations-and-users-in-2025-and-beyond>.
- ²³ Gabriel Morgan Asaftei et al., "State of AI trust in 2026: shifting to the agentic era," McKinsey & Company, March 25, 2026, <https://www.mckinsey.com/capabilities/tech-and-ai/our-insights/tech-forward/state-of-ai-trust-in-2026-shifting-to-the-agentic-era>.
- ²⁴ Asaftei et al., "State of AI trust in 2026."
- ²⁵ Asaftei et al., "State of AI trust in 2026."
- ²⁶ McKinsey & Company, "The state of AI: how organizations are rewiring to capture value," March 12, 2025, <https://www.mckinsey.com/capabilities/quantumblack/our-insights/the-state-of-ai-how-organizations-are-rewiring-to-capture-value>.
- ²⁷ Joe Atkinson et al. "Want ROI from AI? Go for growth." PwC, April 13, 2026, <https://www.pwc.com/gx/en/1/issues/tech-data-ai/ai-roi.html>.
- ²⁸ Sonnenfeld et al., "Anthropic's most powerful AI model just exposed a crisis in corporate governance."
- ²⁹ Beatrice Nolan, "An AI-powered coding tool wiped out a software company's database, then apologized for a 'catastrophic failure on my part,'" *Fortune*, July 23, 2025, <https://fortune.com/2025/07/23/ai-coding-tool-replit-wiped-database-called-it-a-catastrophic-failure/>.
- ³⁰ IBM, Cost of a data breach report 2025, <https://www.ibm.com/think/x-force/2025-cost-of-a-data-breach-navigating-ai>.
- ³¹ Nicole Gillespie et al., *Trust, attitudes and use of artificial intelligence: a global study 2025*, The University of Melbourne and KPMG, 2025, <https://kpmg.com/xx/en/our-insights/ai-and-technology/trust-attitudes-and-use-of-ai.html>.
- ³² Infocomm Media Development Authority, "Singapore launches new model AI governance framework for agentic AI," press release, January 22, 2026, <https://www.imda.gov.sg/resources/press-releases-factsheets-and-speeches/press-releases/2026/new-model-ai-governance-framework-for-agentic-ai>.
- ³³ NIST, *AI RMF 1.0*
- ³⁴ Anthropic, "How we contain Claude across products," May 25, 2026, <https://www.anthropic.com/engineering/how-we-contain-claude>.
- ³⁵ Infocomm Media Development Authority, "Singapore launches new model AI governance framework for agentic AI."
- ³⁶ Gartner, Gartner predicts that Guardian Agents will capture 10–15% of the agentic AI market by 2030, press release, June 11, 2025, <https://www.gartner.com/en/newsroom/press-releases/2025-06-11-gartner-predicts-that-guardian-agents-will-capture-10-15-percent-of-the-agentic-ai-market-by-2030>.
- ³⁷ Derek Mead, "The beautiful, 60-year evolution of the Formula 1 pit stop," *Vice*, April 14, 2014, <https://www.vice.com/en/article/the-beautiful-60-year-evolution-of-the-formula-1-pit-stop/>.



ServiceNow Futures

Our forward-looking research, thought leadership, and platform builds serve as the blueprint for AI-first business reinvention. We help our colleagues and customers see further so they can act sooner.

servicenow.com/futures

servicenow

© 2026 ServiceNow, Inc. All rights reserved. ServiceNow, the ServiceNow logo, Now, and other ServiceNow marks are trademarks and/or registered trademarks of ServiceNow, Inc. in the United States and/or other countries. Other organization names, product names, and logos may be trademarks of the respective organizations with which they are associated.

Version 1.0