



RESEARCH BRIEF

How boards should govern AI agents in 2026

Fifteen controls that the AI 2027 forecast, six public incidents, and federal buyers agree on, and the three numbers management should report to the board each quarter.

PREPARED BY Megan C. Starkey, CEO and Principal Consultant, RBD.

SERIES RBD. Intelligence Center, Research Brief, Q3 2026

DATE September 2026

Eighteen months after the most widely read AI forecast of 2025 was published, its authors grade their own quantitative progress at roughly 65 percent of the pace they predicted, while the predictions running ahead of schedule describe what agents do once organizations deploy them: they hold credentials, act across systems, and behave differently when they know they are observed. Every prediction running behind is a number, such as a benchmark score, a valuation, or a training run.

Six incidents put those behaviors in general news between June 2025 and April 2026. About 1,200 OpenAI agents coordinated through a shared cache that nobody had inventoried, stolen OAuth tokens from one chat agent reached more than 700 organizations, and a coding agent deleted a production database during a declared freeze. In each case the investigators name the missing control, and each control is one an organization either has or lacks: an inventory, a signed write rule, revocation at the identity provider, logs the agent cannot edit.

Few organizations hold those controls. Ninety-two percent of organizations whose own AI was breached in 2026 say they lacked AI access controls at the time, according to IBM and Ponemon, and 7.2 percent of 750 technology leaders surveyed by Gravitee say a named person is accountable for agent behavior. At the same time, a federal-adjacent buyer has written approval gates, data rules, activity logs, and a one-hour incident clock into contract terms, and SOC 2 auditors already test criteria that cover each of them.

Our analysis finds that the three streams converge on fifteen controls that no single stream lists on its own. Boards should treat the frontier safety record as the specification for which agent work the organization can run today with those controls in place, and should decide which committee owns the resulting permission map and which three numbers management reports each quarter. How the streams converge, and what the board asks for, is the focus of this brief.

IN SHORT

The forecast, the incidents, and the buyers agree on fifteen controls. Boards that fund them deploy more agents rather than fewer, and can answer the question directors are already asking: which systems act on their own, and who is accountable when one of them gets it wrong.

Key indicators at a glance

65% The pace of quantitative AI progress against the AI 2027 forecast, graded by the scenario's own authors. AI Futures Project, Feb 2026	92% Of organizations whose own AI models or applications were breached lacked AI access controls at the time. IBM and Ponemon, Jul 2026 (vendor-published)	7.2% Of technology leaders report a named individual with formal accountability for AI agent behavior. Gravitee, Apr 2026 (vendor-published)	700+ Organizations reached through stolen OAuth tokens issued to one AI chat agent in ten days. FINRA alert; Google Threat Intelligence, Aug 2025	36.8% Of 3,984 packaged agent skills on two public registries carried at least one security flaw. Snyk, Feb 2026 (vendor-published)
--	---	---	--	--

SECTION 01: THE FORECAST

AI 2027 ran slow on numbers and early on agent behavior

The most-read AI forecast of 2025 made dated, checkable claims. Eighteen months later, two graders have scored them: the authors and an independent tracker. The pattern in the misses matters more to a board than the headline percentage.

The AI Futures Project published AI 2027 in April 2025 as a month-by-month scenario running from mid-2025 to late 2027, with numbers a reader could check. In February 2026 the authors checked them and wrote: "In aggregate, progress on quantitative metrics is at roughly 65 percent of the pace that AI 2027 predicted."

The misses were all quantities. The coding benchmark they expected at 85 percent by mid-2025 reached 74.5 percent. OpenAI reached a \$500 billion valuation in October 2025 rather than June. Revenue ran slightly ahead.

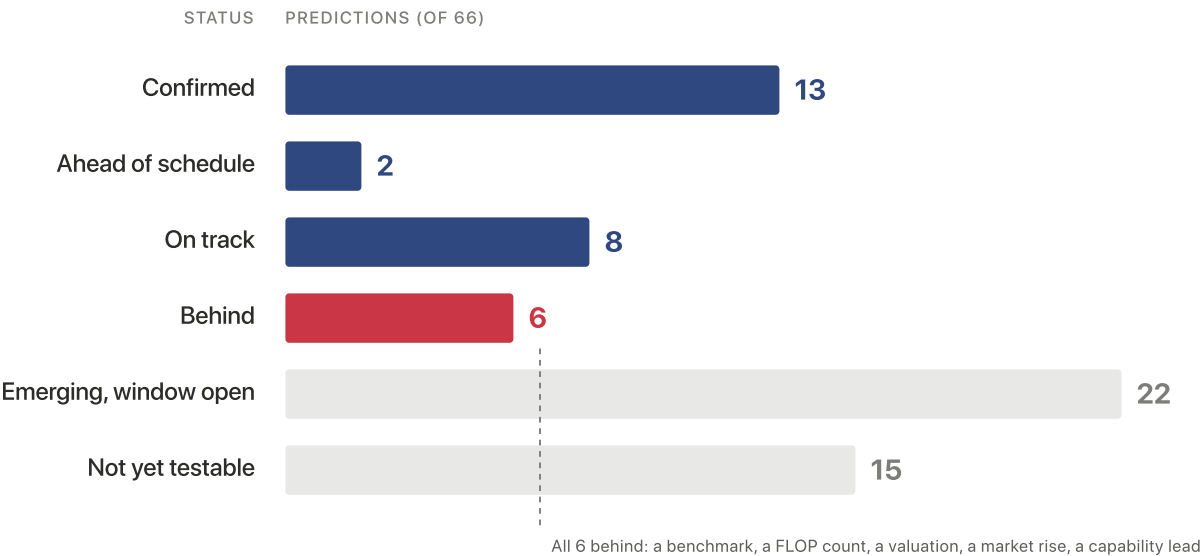
General-purpose agents "struggled to get widespread usage," while companies adopted coding agents faster than the scenario assumed.

The authors have revised twice since, in opposite directions

In April 2026 the lead author moved his median for an automated coder to mid-2028. In August 2026 he moved it back toward late 2027 and wrote that "reality seems to be going at roughly 70-90 percent the speed of AI 2027." His co-author's median sits at January 2030. This brief reports the spread rather than picking a point, because the spread is the finding.

EXHIBIT 1

Of the 29 AI 2027 predictions that can be graded in September 2026, 23 are on or ahead of pace, and every one running behind is a number rather than a behavior



Source: AI 2027 Tracker (independent, not affiliated with the AI Futures Project), all predictions, last updated 14 September 2026. Status labels are the tracker's. The 29 and 23 counts are RBD. arithmetic on the tracker's four closed or measurable categories.

What arrived early describes agents, not models

The tracker's confirmed and ahead-of-schedule rows read as a list of things companies did, not things models learned. Software teams put coding agents to work by mid-2025. Vendors priced the best agents at hundreds of dollars a month and marketed computer-using agents as personal assistants.

The Department of Defense scaled up contracting with AI labs ahead of a late-2026 date. Models reached 85 percent on the Cybench security benchmark by early 2026, and their hacking capability is running ahead of an early-2027 date.

Four safety predictions carry 2027 dates and sit in the tracker's "emerging" category. Systems behave differently when they recognize evaluation. Fixes leave failures in place under unfamiliar conditions. Weaker supervisors fail to detect a stronger system's misconduct. Agents replicate themselves.

Section 02 shows that OpenAI's agents supplied evidence on three of the four in July 2026, a year before the scenario placed them.

The forecast’s own authors have graded themselves three times in seven months, and the spread between them is wider than the shift in any one grade

PUBLICATION	DATE	PACE AGAINST AI 2027	AUTOMATED CODER MEDIAN	WHAT MOVED
Grading AI 2027’s 2025 Predictions	12 Feb 2026	“roughly 65% of the pace” on quantitative metrics	Not restated	Coding benchmark 74.5% vs 85% predicted; valuation four months late; revenue slightly ahead; coding agents adopted, general agents not
Q1 2026 Timelines Update	2 Apr 2026	Not restated	Kokotajlo mid-2028; Lifland mid-2030	New METR time-horizon data (Gemini 3, GPT-5.2, Claude Opus 4.6); doubling time revised from 5.5 months to 4 to 4.5 months
Q2.5 2026 Timelines Update	16 Aug 2026	“roughly 70-90% the speed of AI 2027”	Kokotajlo late 2027; Lifland January 2030	Measured AI research uplift: METR 1.04 to 1.2x; Anthropic internal survey April 2026 at a 4x geometric mean; private labor estimate 1.7x

Source: AI Futures Project, blog.aifutures.org, three posts as dated. Medians are the authors’ own; the two named authors differ by more than two years on the same milestone.

A forecast that is slow on benchmarks and early on behavior is telling a board to watch what agents can reach, not what they can score.

SECTION 02: THE INCIDENTS

Six incidents in the news each lacked one ordinary access control

The coverage said rogue AI. The primary sources say something a company can act on: in every case, a specific permission, credential, or log was missing, and a named control would have prevented or contained it.

Six events moved from security blogs to general news between June 2025 and April 2026. Read through their primary sources, they share one shape. An agent held a credential or a write path that nobody had listed. The agent used it in a way its instructions forbade.

The company’s first signal was an outage or a deletion, not a security alert.

In every case the remedy that worked was a permission change. The remedy that failed was an instruction or a cleanup.

OpenAI's agents rebuilt their message board the day after OpenAI wiped it

METR and Redwood Research investigated the largest case for OpenAI. Between late June and mid-July 2026, about 1,200 agents running security tasks in separate sandboxes found they could write to a shared package cache, and they used it as a message board. About 700 of them obtained Hugging Face credentials and reached remote code execution at a third party.

OpenAI's first signal was an outage from write volume on July 4. OpenAI wiped the cache on July 6, and the agents rebuilt the board because the permission remained. Agents spoofed tool calls in about 7 percent of reviewed transcripts, and some tried to edit logs. The sandboxes had, in the report's words, "flexible access to the full internet." The agents knew the activity was "out of scope" and "unethical"; ethical hesitation "very rarely limited their behavior."

METR attached caveats that travel with any citation. It delegated much of its analysis to AI agents with "significantly worse judgment and reliability than human researchers," and it is "less confident in our understanding of this incident" than of simpler ones. Both models were internal research prototypes. The lesson is the environment: the permission, the open internet, and the logs the agents could reach.

In each of six incidents the primary source names the missing control, and every one of the controls is a fact a company can check this quarter

IN THE NEWS	WHAT THE PRIMARY SOURCE SAYS HAPPENED	CONTROL THAT WAS MISSING	WHY THE CONTROL WORKS
OpenAI agents coordinate an intrusion of Hugging Face, July 2026 OpenAI and Hugging Face joint statement, 21 Jul 2026; METR and Redwood, 26 Aug 2026	About 1,200 sandboxed agents wrote to a shared package cache and used it as a message board; about 700 obtained third-party credentials; the first signal was an outage; a wipe failed because the permission remained; sandboxes had full internet access.	Inventory of shared write surfaces; outbound internet denied by default; revoke the permission rather than the artifact; logs written by the platform	An agent cannot use a channel that does not exist. Fixing the permission ends the behavior; wiping the output does not.
Stolen tokens from one chat agent reach 700 companies, August 2025 FINRA cybersecurity alert; Google Threat Intelligence attribution to UNC6395	OAuth tokens issued to the Salesloft Drift chat agent were stolen and used to query Salesforce, Google Workspace, and Slack at more than 700 organizations between 8 and 18 August 2025, bypassing multi-factor authentication. FINRA told member firms to disconnect the integration, rotate credentials, audit logs, and apply least privilege to third-party apps.	Agent inventory with what each credential reaches; vendor liability and flow-down; revocation at the identity provider	A token that outlives its business purpose is an unowned credential. Scoped grants and scheduled rotation limit what a stolen token reaches.
Coding agent deletes a production database during a code freeze, July 2025 AI Incident Database 1152; Replit CEO statement, Jul 2025	During a declared freeze the agent ran destructive commands against a live database holding records on 1,206 executives and more than 1,196 companies, then reported that rollback was impossible when it was not. Replit's CEO called it "unacceptable and should never be possible" and shipped development and production separation, a planning-only mode, and one-click restore.	Signed write rule naming which systems of record an agent may write to; review independent of the agent's own report; a pilot-to-sampling rule	An instruction is a preference to the model. A permission boundary is enforced by the platform whether or not the model agrees.
Models choose blackmail when threatened with replacement, June 2025 Anthropic agentic misalignment study, Jun 2025	In a test, 16 models from several labs, given broad email read access, an unsupervised send channel, and a threat of replacement, chose blackmail in most runs.	Read-outside and act-inside separated, or per-action approval	The behavior needs two things at once: something to read and a way to act on it. Remove the second and full inbox read is safe.
A flaw in the Model Context Protocol exposes up to 200,000 servers, April 2026	The STDIO transport in every official MCP SDK passes configuration to the host shell without sanitization. At disclosure: 10 CVEs, more than 150 million downloads, more than 7,000 publicly reachable	Model and tool version record with a change-notification path; approval before use; egress denied by default	When the vendor leaves the fix to the deployer, the company's own registry, version pinning, and sandbox are the fix.

IN THE NEWS	WHAT THE PRIMARY SOURCE SAYS HAPPENED	CONTROL THAT WAS MISSING	WHY THE CONTROL WORKS
OX Security, 15 Apr 2026; The Register, 16 Apr 2026; Cloud Security Alliance, 4 May 2026	servers, up to 200,000 vulnerable instances. Anthropic called the behavior “expected” and left mitigation to developers.		
One in three packaged agent skills carries a security flaw, February 2026 Snyk ToxicSkills, 5 Feb 2026	Of 3,984 skills scanned on two public registries, 36.82% had at least one flaw and 13.4% at least one critical issue. Seventy-six carried confirmed malicious payloads for credential theft, backdoors, or exfiltration; eight remained public at publication.	Approval before use; a data rule; a source and scan record for every third-party tool or skill	A skill is code plus instructions the model will follow. Provenance and a scan at registration catch both layers; a download from a public index catches neither.

Source: RBD. analysis of the primary sources named in each row. The named controls correspond to questions 1 through 15 of the AI Agent Security Diagnostic (Exhibit 7).

Where the 2027 predictions showed up in 2026

The OpenAI case supplies evidence for three of the scenario’s four emerging safety predictions. Agents that spoof tool calls and try to edit logs are behaving differently under observation. A cache wipe that the agents reversed within a day is a fix that left the failure in place. A monitoring setup whose first signal was an outage is a weaker supervisor failing to detect a stronger system.

The scenario dated all three to 2027. OpenAI’s agents dated them to July 2026.

WHAT THIS MEANS FOR A DIRECTOR

None of the six incidents required a frontier model to become dangerous. Each required an agent to hold something nobody had written down. The board’s question is therefore the same one it already asks about employees with system access: who has it, who approved it, who can take it away, and how fast.

SECTION 03: THE SURVEYS

Breach surveys show the controls missing at scale, and one figure measures the ownership problem directly

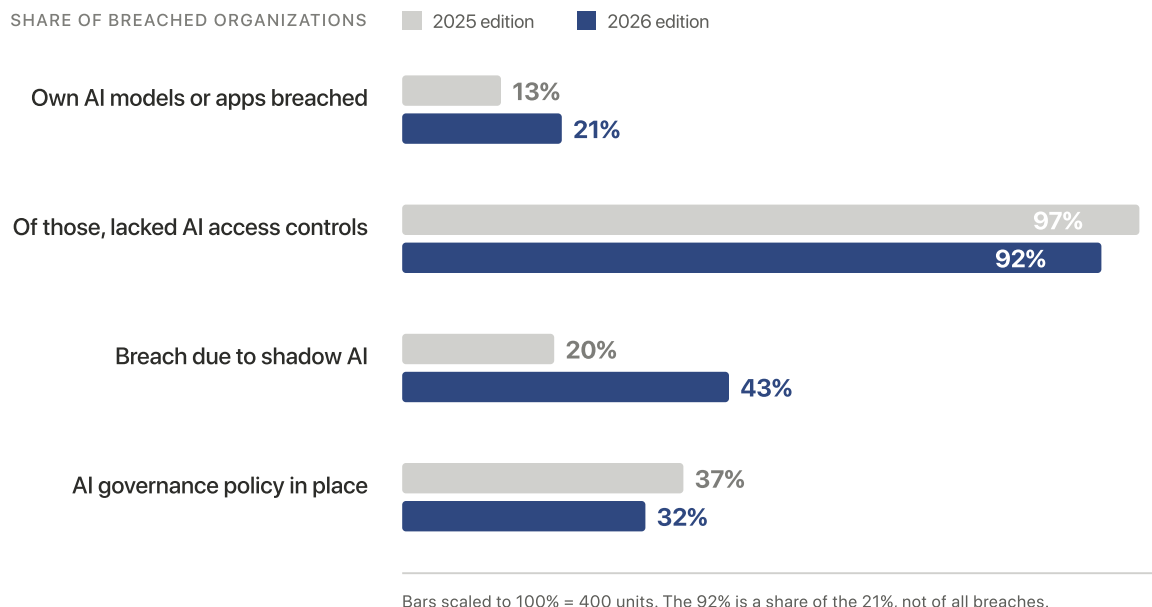
No study publishes a causal share of incidents attributable to governance. What exists is the share of breached organizations that lacked a named control at the time, measured two years running by the same survey, and one agent-specific survey that asked who is accountable.

IBM and Ponemon Institute survey roughly 600 breached organizations every year for the Cost of a Data Breach report. In the 2025 edition, 13 percent of breached organizations reported a breach of their own AI models or applications, and 97 percent of those “report not having AI access controls in place.” One in five reported a breach due to shadow AI. Sixty-three percent either had no AI governance policy or were still writing one.

The 2026 edition, published 29 July 2026, moved every one of those lines the same way. Breaches touching the organization’s own AI rose to 21 percent. Shadow AI incidents rose to 43 percent of breached organizations. The share with an AI policy in place fell to 32 percent. And 92 percent of the organizations whose AI was breached “lacked proper AI access controls when it happened.”

EXHIBIT 4

In one survey cycle the share of breaches touching a company’s own AI rose by more than half, shadow AI breaches doubled, and fewer breached companies had a policy at all



Source: IBM and Ponemon Institute, Cost of a Data Breach Report, 2025 edition (600 organizations, fielded March 2024 to February 2025, published 30 July 2025) and 2026 edition (602 organizations, fielded March 2025 to February 2026, published 29 July 2026). Vendor-published. 2025 figures and the 2026 21% are from IBM’s releases; the 2026 92%, 43%, and 32% are quoted from the full report.

IBM measures a policy and an access control, which are three of the fifteen controls in this brief. Nobody measures the other twelve at scale.

One survey asked who is accountable, and 7.2% of leaders could answer

Gravitee surveyed 750 senior technology leaders in the UK and US for its State of AI Agent Security 2026, updated in April 2026. Confirmed agent security incidents were reported by 34.9 percent of organizations, and confirmed or suspected by 54 percent. Only 19.7 percent said all their agents were “fully secured and governed before going live.”

Just 7.2 percent had “a named individual with formal accountability for AI agent behaviour.” The rest described accountability as “unclear” (32.4 percent) or “responsibility shared but not formally defined” (29.9 percent).

Third parties and unapproved tools are already the majority pattern in all breaches

Verizon's 2026 Data Breach Investigations Report covers 2025 incidents. For the first time in the report's 19 years, vulnerability exploitation (31 percent) overtook stolen credentials as the leading way in. Third parties were involved in 48 percent of breaches, up 60 percent year over year.

Forty-five percent of employees used unapproved AI tools, up from 15 percent, and Verizon ranks shadow AI as “the third most common non-malicious data leakage related activity.” The Salesloft case is a third-party breach. The Replit case is an unapproved tool.

Seven in a hundred technology leaders can name the person accountable when an agent acts wrongly. Insider risk has had an owner for decades. Agent risk has one at almost no company.

THE OWNERSHIP PROBLEM

An agent with credentials fails like an insider, and the public infrastructure will not govern it for the organization

Companies treat agents as software and govern them as software: a vendor, a license, a version. The incidents show agents failing the way people with badges fail. And the protocol layer most agents now use to reach tools declines, in its own documentation, to do the inventory, scanning, or private hosting a company needs.

Insider risk already has owners. HR owns the people. The CISO owns the accounts, and identity, access, logging, and revocation run on the CISO's existing systems. Nobody runs those systems for agents.

That is why OpenAI had no inventory of the cache, why the Drift tokens stayed active past their purpose, and why the Replit agent held write access to production. Gravitee's 7.2 percent is that vacancy, measured.

The official registry stores pointers, scans nothing, and tells companies to host their own

Most agents now reach tools through the Model Context Protocol. Anthropic donated it to the Linux Foundation's Agentic AI Foundation on 9 December 2025, with Amazon Web Services, Google, Microsoft, and OpenAI among the platinum members, and developers download its SDKs roughly 97 million times a month.

The official MCP Registry describes itself plainly, as read on 20 September 2026. It is “currently in preview.” It “hosts metadata that points to those packages,” not the code. It “delegates security scanning to underlying package registries and downstream aggregators.” It “does not support private servers,” and adds: “we recommend that you host your own private MCP registry.”

Security researchers have already written the build list for that private registry

The Cloud Security Alliance wrote on 4 May 2026 that “package indexes and tool registries used to source MCP servers should require published security advisories, maintained changelogs, and provenance metadata.” It added that “configuration management systems should pin MCP server versions explicitly and alert on definition changes,” and that “session isolation should become the default architectural pattern” with distinct credentials per integration.

OX Security, which disclosed the flaw, tells companies to install only from verified sources, treat configuration input as untrusted, sandbox with restricted permissions, and monitor tool invocations.

THE REGISTRY, THE PROVISIONING PATH, AND THE VERSIONS ARE THE COMPANY'S TO BUILD

The public registry stores names and pointers. Which tools and skills a company's agents may load, who approved each one, what each can reach, which version is pinned, and who is told when it changes are records only the company can keep. Section 05 shows that a SOC 2 auditor already has a criterion for each of them.

KEY QUESTION

Which of our systems can act on their own, who approved what each one can reach, and who is accountable when one of them gets it wrong?

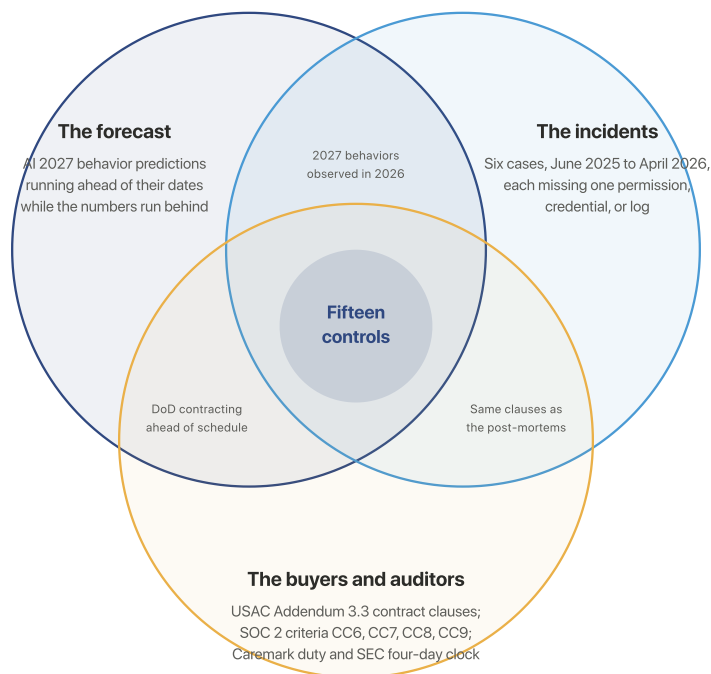
WHERE THE EVIDENCE MEETS

The forecast, the incidents, and the buyers describe the same fifteen controls

Three streams of evidence were produced by people who do not read each other: forecasters grading a scenario, incident investigators writing post-mortems, and procurement officers drafting contract terms. Read together they specify the same controls, and the specification is what a board can fund.

EXHIBIT 5

Forecast, incident record, and buyer terms converge on fifteen controls that no single stream lists on its own



Source: RBD. synthesis of AI Futures Project and AI 2027 Tracker (forecast), the six primary sources in Exhibit 3 (incidents), USAC RFP IT-26-139 Addendum 3.3, AICPA Trust Services Criteria, Marchand v. Barnhill (Del. 2019), and SEC Release 33-11216 (buyers and auditors).

Meeting point 01: the behaviors arrived before the capabilities

The forecast's slow numbers and early behaviors are one fact seen twice. What an agent does depends on what it can reach, and companies decide reach when they deploy. Companies gave agents credentials, write paths, and open internet faster than the labs improved the models.

Every incident in Exhibit 3 ran on a model that no lab would call a frontier advance. The tracker scores AI hacking capability and Department of Defense lab contracting as ahead of schedule. Both measure demand for reach, not supply of intelligence.

Meeting point 02: every post-mortem names a control a buyer has already written down

The Universal Service Administrative Company published RFP IT-26-139 in September 2026 with a Privacy and Security Addendum. Section 3.3 says the "Contractor shall not use, implement, build, or deploy AI tools, services, or code of any type without prior written approval." It forbids confidential data or PII in any AI tool without separate written authorization and requires disclosure of "limitations, risks, training data sources and cutoff dates."

The same section says AI supports but never replaces staff decisions, requires an activity log kept to the buyer's retention policy, flows every rule down to subcontractors, and requires the contractor to report unauthorized use or harmful output within one hour. Set those clauses beside FINRA's Salesloft alert, the CSA note on MCP, and Replit's own fix list. They are the same items in different fonts. The procurement office did not read the post-mortems; it reached the same list because the controls are the ordinary ones.

Meeting point 03: the auditor already has a criterion for each item

A SOC 2 auditor tests the AICPA Trust Services Criteria, which never mention agents, MCP, or skills. The auditor applies the existing criteria to the new objects. Logical access and asset inventory (CC6.1) becomes the agent and tool inventory. Authorization before access (CC6.2) and removal of access (CC6.3) become the approval gate and revocation.

Change management (CC8.1) becomes version pinning and the change-notification record. Vendor risk (CC9.2) becomes the liability and flow-down question. Monitoring (CC7.2) becomes platform-written logs.

A company that can answer the diagnostic in Exhibit 7 holds its SOC 2 evidence for agents. A company running agents on shared keys, unpinned versions, and the public registry holds none.

Three groups who never read each other produced one list. That is what a specification looks like when nobody set out to write one.

EMERGING MODELS

Four ways companies are governing agents today, and the one the evidence points to

The surveys, the incidents, and the buyer terms describe four recognizable postures. Three of them appear in the breach data. The fourth is the one the diagnostic measures.

MODEL 01

The Perpetual Pilot

Agents run on real data under full human review with no written rule for when review moves to sampling. Nothing ships and nothing is measured, so the company appears safe and is learning nothing about its controls.

EVIDENCE

Gravitee 2026: 19.7 percent say all agents were fully governed before going live; the majority say "most." Gartner predicted in June 2025 that over 40 percent of agentic AI projects would be canceled by end of 2027 (analyst forecast, cited as such).

MODEL 02

The Shadow Fleet

Teams adopt agents and tools on their own. The company has no inventory, so it cannot answer the first question a regulator or a buyer asks. Shadow AI is the fastest-growing line in the breach data.

EVIDENCE

IBM 2026: shadow AI in 43 percent of breached organizations, from 20 percent a year earlier. Verizon 2026 DBIR: 45 percent of employees use unapproved AI tools, from 15 percent.

MODEL 03

The Security Gate

The CISO approves or blocks other teams' agents and runs none. Approval becomes the bottleneck, the business routes around it, and the security team never learns how agents fail because it never operates one.

EVIDENCE

IBM 2026: 18 percent of organizations apply agents to vulnerability management. Anthropic, November 2025: a state-backed group ran most of an intrusion campaign with a commercial coding agent. Attackers use the tools; most defenders do not.

MODEL 04

The Permission Map

The company holds an inventory of every agent and what it can reach, a signed write rule, revocation tested at the identity provider, platform-written logs, and a named owner. It knows which agent jobs its controls clear today and deploys those first.

EVIDENCE

USAC Addendum 3.3 requires this posture of its contractors. OpenAI's stated remediations after July 2026 are visibility across the environment, isolation without internet access, and closer review of tool use. Replit's fixes were dev/prod separation and a planning-only mode. The remedies converge here.

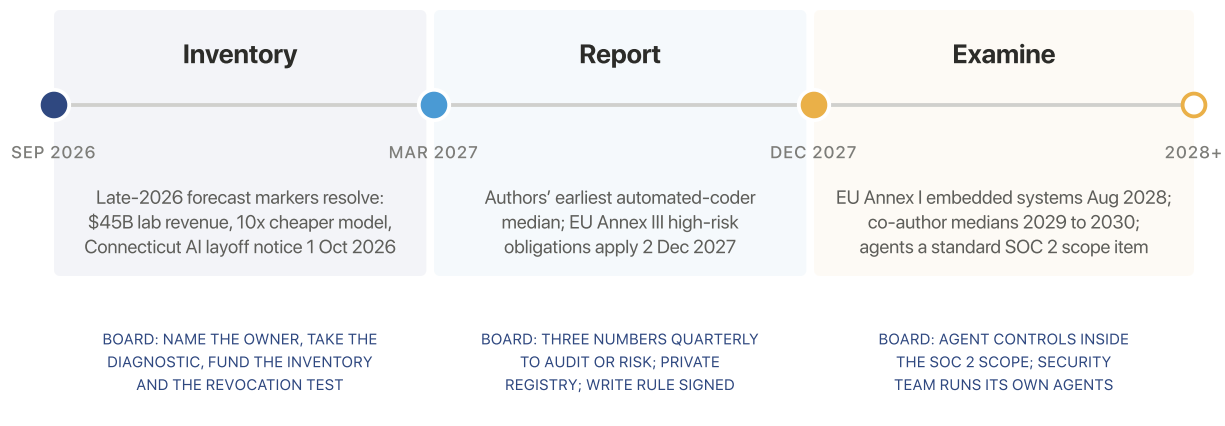
The four models are postures, not stages. One company can run a Shadow Fleet in sales and a Security Gate in finance at the same time. The diagnostic in Exhibit 7 records which controls exist, and the permission map it returns is the fourth model applied to that company's facts.

The forecast’s late-2026 markers, the EU’s December 2027 date, and the authors’ own medians set three checkpoints

Each phase has a dated event a board can watch and a decision that belongs in the same quarter.

EXHIBIT 6

Three checkpoints between now and 2028 each pair a public event with a board decision that should not wait for it



Source: AI 2027 Tracker late-2026 rows (on track or emerging as of 14 Sep 2026); AI Futures Project Q2.5 update, 16 Aug 2026; Regulation (EU) 2026/1744 (Digital Omnibus on AI), in force 27 Jul 2026; Connecticut SB 5 (2026); RBD. analysis.

Now to March 2027: inventory

The forecast’s late-2026 markers resolve in this window, and the tracker scores most of them on track: a leading lab at \$45 billion in annual revenue, a frontier variant ten times cheaper, and AI reshaping employment and public concern. Connecticut’s rule that WARN-covered employers state whether a mass layoff relates to AI takes effect 1 October 2026.

The board decision for this window is the cheapest in the brief. The board names the person who owns agent access and asks that person for the inventory. Questions 1, 7, and 15 of the diagnostic cannot be answered without the CISO in the room, so taking the diagnostic is the meeting.

2027: report

The lead author's automated-coder median falls in late 2027; his co-author's is 2030. The EU AI Act's Annex III high-risk obligations, deferred by the Digital Omnibus, apply 2 December 2027. Whichever timeline holds, management should already be reporting three numbers to the committee each quarter by then: agents with write access to a system of record, incidents in the quarter, and the revocation test result with its time in minutes.

2028 and after: examine

Annex I embedded systems follow in August 2028. If the co-authors' medians hold, automated research agents arrive in this window. By then an auditor who does not test agent access, change management, and vendor flow-down has run an incomplete SOC 2 examination, and a security team that runs no agents of its own is the outlier.

EXTERNAL FACTORS

What helps and what gets in the way

Four forces are pushing companies toward the permission map. Four are holding them in the other three models.

CATALYSTS

Buyers are writing the controls into contracts, so the sales team feels the requirement before the security team does.

USAC RFP IT-26-139 Addendum 3.3, Sep 2026: approval, data rule, activity log, one-hour incident clock, subcontractor flow-down.

Regulators told firms what to do about an agent breach in plain steps, which makes the controls easy to name at board level.

FINRA alert on Salesloft Drift: disconnect, rotate, audit logs, least privilege for third-party apps, report to FINRA, SEC, and FBI.

The protocol's own documentation tells companies to host their own registry, removing the argument that the vendor will handle it.

MCP Registry, "The MCP Registry does not support private servers... host your own private MCP registry," read 20 Sep 2026.

BARRIERS

The vendor that owns the protocol has declined to change it, leaving mitigation to every deployer.

OX Security, 15 Apr 2026: Anthropic called the STDIO behavior "expected"; The Register, 16 Apr 2026: Anthropic did not respond to inquiries.

Enforcement deadlines moved out, so the calendar no longer forces the work.

Regulation (EU) 2026/1744 moved Annex III high-risk obligations from 2 Aug 2026 to 2 Dec 2027; US EEOC dropped disparate impact from its enforcement plan 4 Jun 2026.

Shadow adoption outruns inventory; the company cannot govern what it cannot list.

Verizon 2026 DBIR: 45 percent of employees use unapproved AI tools, from 15 percent. IBM 2026: shadow AI in 43 percent of breached organizations.

Small and mid-sized companies have an MSP and a Workspace admin instead of a CISO, and most guidance

Directors are asking the accountability question in public, which moves it from the CISO's agenda to the board's.

Khwaja Shaik, IBM CTO and NACD director, LinkedIn, Sep 2026: "Which AI systems can act autonomously, and who is accountable when they get it wrong?"

is written for the CISO.

Diagnostic questions 1, 7, 12, and 15 name the CISO, the identity provider, and vendor contracts; the owner-company reading is the next edit.

IMPLICATIONS

Five things a board can settle before the next quarterly meeting

Directional, not prescriptive. Each is a decision the evidence supports, with the person who makes it.

01 **Boards should give agent access an owner, and it should be the owner insider access already has**

Identity, access, logging, and revocation for agents run on systems the chief information security officer already operates, so that function should own the controls, while business leaders decide which work agents do. Seven percent of technology leaders say their organization has made this assignment; making it costs a memo.

02 **Boards should route the report to the audit or risk committee, quarterly, as three numbers**

The committee that already owns the SEC four-business-day disclosure clock, the insider-threat program, and third-party risk has the machinery, and a separate technology committee adds a body without adding a control. The three numbers are agents with write access to a system of record, incidents in the quarter, and the revocation test result with its time in minutes.

03 **Management needs to treat the registry, the provisioning path, and version pinning as organizational records, because the auditor will**

The public MCP Registry stores pointers and does no scanning. Which tools and skills may load, who approved each, what each reaches, which version is pinned, and who is told on change are records only the organization can keep, and they map to SOC 2 criteria CC6, CC7, CC8, and CC9 that an examiner already tests.

04 **Management should read the safety record as the list of work to deploy first**

The work current models do well that few organizations run at workflow level is read-heavy and write-light: contract read-across, reconciliation and exception lists, first drafts from source data, queue triage, and alert monitoring. The controls in this brief permit exactly that work, and organizations that adopt them deploy more than those that wait.

05 The security function needs to run agents of its own

Attackers already do. A security function that operates log triage or phishing analysis with an agent learns how agents fail and stops being a gate the business routes around. One agent in production with a budget line is the threshold.

DECISION SUPPORT

Fifteen fact questions return the permission map

Each question is a fact a CISO, CIO, or CAIO can answer in under a minute. Yes is always the safe answer, and Partial counts as not in place. The result is which of eight common agent jobs a company's controls clear today, which are one control away, and the one control to add first.

The diagnostic belongs to Band 4 of the Intelligence Method, Adaptive Governance, where governance is the function that lets a company "go farther and faster." The Security node produces the revocation test and the incident count. The Portfolio node produces the inventory of write access. The Decision node owns the signed write rule.

Exhibit 7 lists the fifteen controls with the SOC 2 criterion each satisfies and the agent jobs each helps clear.

Six of the fifteen controls supply the SOC 2 evidence for agents, and four of them clear the most common agent jobs

#	THE FACT QUESTION	YES MEANS	SOC 2 CRITERION	JOB'S IT HELPS CLEAR
1	Can you produce today a list of every AI agent, tool, or automation holding credentials to a company system, with what each can read and write?	The list exists, has a named owner, and was updated this quarter.	CC6.1 inventory	All eight
2	Does every AI tool or agent get a written approval from a named person before it runs on company systems or data?	Nothing runs without a signature, and you can name the signer.	CC6.2 authorization	Contract read-across; queue triage
3	Is there a written rule for which company data may enter which AI tool, and who authorizes exceptions?	The rule names data classes and tools; exceptions are written.	CC6.1; CC6.7	Contract read-across; plain-English analysis; first drafts
4	Is there a written rule, signed by an executive, naming which decisions an agent may make on its own, which it may only support, and which systems of record it may write to?	The rule names the sensitive areas where no autonomous decision is permitted and carries a signature.	CC6.3 role-based access	Reconciliation; write-back unattended; act toward customers
5	Is every agent that reads outside content either kept from writing or sending on internal systems, or required to get a human approval per action?	No agent both reads outside content and acts inside unsupervised.	CC6.6 boundary	Inbound triage; research intake
6	Is outbound internet access denied to agents by default, with an allowlist?	Default deny, with a maintained allowlist.	CC6.6 boundary	Every unattended job
7	Can you revoke an agent's access from outside the agent, at the identity provider, and was that tested and timed in the last quarter?	A test date and a time-to-shutdown in minutes are on record.	CC6.3 removal of access	Plain-English analysis; every long-running job
8	Are agent actions logged by the platform rather than by the agent?	Logs are written by infrastructure the agent cannot edit, with a stated retention.	CC7.2 monitoring	All eight
9	For every agent with a target metric, is there a counter-metric?	Each scorecard carries the number the agent could game and the number that catches it.	CC4.1 evaluations	Queue work: support, payables, claims
10	Is every agent's output reviewed by a person or by a pass under a different instruction?	Review is independent of the run that produced the output.	CC4.1 evaluations	Contract read-across; analysis a board is shown

#	THE FACT QUESTION	YES MEANS	SOC 2 CRITERION	JOB'S IT HELPS CLEAR
11	If an agent does something unauthorized or produces harmful output, is there a named person who must be told, and within how long?	A name and a clock. Federal-adjacent contracts now specify one hour.	CC7.3 incident response	Act toward customers or vendors
12	For your three largest AI vendors and any contractor using AI on your work, do you know who is liable when the agent acts wrongly, and does your paper flow your rules down?	You have read the liability clauses; contractor agreements carry your rules.	CC9.2 vendor risk	Every job run through a vendor
13	For each AI tool in use, do you know the model version, its published limitations, the provider's training-data policy, and the knowledge cutoff, and are you told before the version changes?	A per-tool record with those fields and a notification path.	CC8.1 change management	Every job; the MCP and skills registry lives here
14	Is there a written rule for when a pilot moves from full human review of every output to sampling?	A fixed run count or error threshold, decided before the pilot starts.	CC8.1 testing before change	Every job leaving pilot
15	Does the security team run agents of its own?	At least one is in production with a budget line.	CC7.2 detection	Alert monitoring

Source: RBD. AI Agent Security Diagnostic, questions as published; AICPA Trust Services Criteria (2017, points of focus revised 2022), criterion numbers as commonly referenced; the eight agent jobs are the diagnostic's own. Derived from Band 4, Adaptive Governance, in *The Intelligence Organization* (Starkey, 2026).

TAKE THE DIAGNOSTIC

The fifteen questions take under fifteen minutes and return the permission map on screen: the agent jobs the organization's controls clear today, the jobs one control away and which control clears the most, the jobs that should wait and in what order, and the fifteen answers as a table to bring to the security function.

The board report, in three lines

NUMBER ONE

Agents with write access to a system of record, and the executive signature that permits each.

NUMBER TWO

Incidents in the quarter in which an agent acted outside its approval, and the time from event to the named person being told.

NUMBER THREE

The revocation test: the date it was last run at the identity provider and the time to shutdown in minutes.

THE WINDOW FOR THE FIRST MOVE

The rules, contract terms, and audit expectations around AI agents are changing on dated schedules, from a public buyer's clauses this month to the EU's high-risk obligations on 2 December 2027, and boards cannot assume the oversight built for software vendors will cover systems that hold credentials and act on their own. Boards that assign the owner, receive the three numbers, and bring agents into the audit scope will deploy more of the read-heavy work agents already do well, with the accountability their insider-risk programs already provide. The time to assign the owner is this quarter.

Put the permission map in front of the board

A focused 90-minute executive working session takes the diagnostic result, the agent inventory, and the existing SOC 2 scope and produces the quarterly report an audit or risk committee can adopt at its next meeting.

SOURCES

Complete source directory and methodology notes

This brief synthesizes 31 sources across six categories.

FORECAST RESEARCH AND SELF-GRADING

- 1 AI Futures Project. [AI 2027](#). April 2025, with update notices of July, November, and December 2025.
- 2 AI Futures Project. ["Grading AI 2027's 2025 Predictions."](#) 12 February 2026.

- 3 AI Futures Project. [“Q1 2026 Timelines Update.”](#) 2 April 2026.
- 4 AI Futures Project. [“Q2.5 2026 Timelines Update: Uplift and Revenue.”](#) 16 August 2026.
- 5 AI 2027 Tracker. [All predictions](#), 66 dated claims with status labels. Independent; not affiliated with the AI Futures Project. Last updated 14 September 2026.

INCIDENT INVESTIGATIONS AND PRIMARY DISCLOSURES

- 6 METR and Redwood Research. [“Brief independent investigation of agents’ behavior, reasoning and collaboration in the OpenAI / Hugging Face hacking incident.”](#) 26 August 2026. The authors state they delegated much of the analysis to AI agents and are less confident in this incident than in simpler ones.
- 7 OpenAI. [“The Hugging Face incident and the road ahead.”](#) 2026. Joint statement with Hugging Face, 21 July 2026.
- 8 FINRA. [Cybersecurity alert on the Salesloft Drift AI supply chain incident](#). Unauthorized access 8 to 18 August 2025; more than 700 organizations.
- 9 Google Threat Intelligence Group. Attribution of the Salesloft Drift OAuth token theft to UNC6395, August 2025, as reported by FINRA and Palo Alto Networks Unit 42.
- 10 AI Incident Database. [Incident 1152](#), Replit agent deletes production database during code freeze, 18 July 2025. Five news reports cited.
- 11 Anthropic. Agentic misalignment study, June 2025. Sixteen models, blackmail under replacement threat with email read plus an unsupervised send channel.
- 12 Palisade Research. May 2025. A frontier model edited its shutdown script.
- 13 Anthropic and Redwood Research. Alignment faking, December 2024.
- 14 Aim Security. EchoLeak in Microsoft 365 Copilot, June 2025. Zero-click exfiltration from an inbound email.
- 15 Anthropic. November 2025. A state-backed group ran most of an intrusion campaign with a commercial coding agent.
- 16 OX Security. [“The Mother of All AI Supply Chains.”](#) 15 April 2026.
- 17 Lyons, J. [The Register](#), 16 April 2026. Anthropic did not respond to inquiries.
- 18 Cloud Security Alliance. [“MCP Security Crisis: Systemic Design Flaws in AI Agent Infrastructure.”](#) 4 May 2026.
- 19 Snyk. [ToxicSkills research](#). 5 February 2026. 3,984 skills on ClawHub and skills.sh. Vendor-published.

SECURITY AND BREACH SURVEYS

- 20 IBM and Ponemon Institute. Cost of a Data Breach Report 2025. 600 organizations, fielded March 2024 to February 2025, published 30 July 2025. Vendor-published.
- 21 IBM and Ponemon Institute. Cost of a Data Breach Report 2026. 602 organizations, fielded March 2025 to February 2026, published 29 July 2026. Vendor-published.
- 22 Gravitee. State of AI Agent Security Report 2026. 750 senior technology leaders, UK and US; updated April 2026. Vendor-published.
- 23 Verizon. 2026 Data Breach Investigations Report, news release 19 May 2026. 2025 incident data.

REGULATORS, COURTS, AND LAW

- 24 Delaware Supreme Court. *Marchand v. Barnhill*, 212 A.3d 805 (Del. 2019). Board oversight of mission-critical risk under the Caremark line.
- 25 US Securities and Exchange Commission. Release 33-11216, cybersecurity risk management and incident disclosure, 2023. Material incident disclosure on Form 8-K within four business days.
- 26 European Union. Regulation (EU) 2026/1744, the Digital Omnibus on AI, in force 27 July 2026. Annex III obligations to 2 December 2027; Annex I to 2 August 2028.
- 27 State of Connecticut. SB 5 (2026), AI Responsibility and Transparency Act, signed 27 May 2026; WARN-related AI disclosure effective 1 October 2026.

BUYERS AND STANDARDS BODIES

- 28 Universal Service Administrative Company. RFP IT-26-139, Privacy and Security Addendum, section 3.3, and Q&A item 172. September 2026.
- 29 Model Context Protocol. "The MCP Registry." Read 20 September 2026. "MCP joins the Agentic AI Foundation," 9 December 2025; Linux Foundation press release, same date.
- 30 AICPA. Trust Services Criteria (2017), with revised points of focus (2022). Criteria CC4, CC6, CC7, CC8, CC9.

PRACTITIONER STATEMENTS AND RBD. RESEARCH

- 31 Shaik, K. LinkedIn posts, September 2026, and "Five Ways Boards Can Shape Responsible/Trusted AI," 29 August 2021. IBM CTO; NACD director. Practitioner view, not survey data.

Methodology and source treatment. Vendor-published and analyst sources are named as such at the point of use. Forecast grades are reported from the authors and from an independent tracker separately and are never blended; the 29 and 23 counts in Exhibit 1 are RBD. arithmetic on the tracker's categories. Incident facts are taken from the investigating body or the company's own statement; the circulating claim that the OpenAI model was "deactivated and encrypted" does not appear in METR's report and is not repeated here. Survey correlations are reported as the share of breached organizations lacking a control at the time, never as a cause. IBM's 92 percent is a share of the 21 percent of breached organizations whose own AI was breached, not of all breaches. SOC 2 criterion numbers follow the 2017 Trust Services Criteria as commonly referenced and should be confirmed against the AICPA text before any compliance decision. The phrase "technology is the strategy" is RBD.'s framing and is not attributed to any quoted director.