

DAILY MARKET PULSE · MONDAY CONTRARIAN

Monday, September 21, 2026 · America/New_York

Human-in-the-loop is a posture, not a control

The most useful agent-governance evidence published this year did not come from a regulator. It came from a frontier lab measuring its own traffic, and it says the thing most enterprise AI policies are built on gets weaker exactly where the stakes get higher. Human involvement appears in 87% of agent tool calls on trivial tasks and 67% on complex ones VERIFIED C01. Banking's new model-risk letter, meanwhile, puts agentic AI out of scope entirely CITED C06. This edition maps what regulated operators should measure instead, across financial services, healthcare, manufacturing and energy.

Table of contents

1. The 60-second scan
2. Contrarian read: the approval checkbox
3. Frontier ledger: what actually shipped
4. Robotics reality check: the dexterity numbers
5. Financial services
6. Healthcare
7. Manufacturing
8. Energy and utilities
9. Implementation architecture
10. Scenario planning and risk-reward
11. Practitioner FAQ and did-you-know
12. Sources

87% → 67%

Human involvement in agent tool calls, minimal-complexity tasks vs high-complexity tasks

Anthropic first-party measurement, n = 998,481 tool calls. Anthropic states its human-involvement figures are an upper bound.

VERIFIED C01

0.8%

Share of sampled agent actions that appear irreversible

The small share is the point: reversibility, not approval volume, is what bounds damage.

VERIFIED C01

21%

Organizations reporting a mature governance model for agentic AI

Deloitte survey of 3,235 IT and business leaders across 24 countries.

CITED C11

40% vs 22%

Large enterprises scaling AI agents vs smaller organizations

McKinsey 2026 global survey: large-org scaling rose from 27% a year earlier; smaller organizations were flat.

CITED C10

What moved last week — Monday, September 14 through Friday, September 18, 2026. Anthropic disclosed that Claude now leads roughly 26% of its own research and development work, up from none in February 2026, and participates in over 90% of it, while stating that Claude is not fully autonomous on any measured subset of its work. CITED C02, company-reported SpaceXAI shipped Memory in Grok Build on September 16 and Grok Voice Transcribe 2.0 on September 18. VERIFIED C05 OpenAI published product material on connecting AI usage to business value on September 16 and has DevDay scheduled for September 29, 2026. CITED C04, C22 NVIDIA's latest filing showed supply commitments of roughly \$279 billion, up from roughly \$119 billion three months earlier. CITED C23

THE ARGUMENT IN ONE PARAGRAPH

Enterprises wrote "a human reviews every agent action" into their AI policies because it was the only control that fit in a single sentence. Frontier field data now shows that control decays under exactly the conditions it exists for: long task chains, experienced operators, and high complexity VERIFIED C01.

Regulators are not rescuing it either. The federal banking agencies took generative and agentic AI out of the scope of their rewritten model-risk guidance in April 2026 CITED C06, and the EU deferred standalone high-risk obligations to December 2027 VERIFIED C15. The gap is not a reason to wait. It is a reason to replace an unmeasurable control with four measurable ones: reversibility class, blast radius, time-to-halt, and evidence of intervention.

2. Contrarian read: the approval checkbox is not the control

Almost every enterprise AI policy written in the last eighteen months contains a version of the same clause: a qualified human reviews and approves agent actions before they take effect. It reads well in a board pack. It survives a first-round audit. And the largest published field study of how agents are actually supervised suggests it does not hold where it matters.

Anthropic analyzed 998,481 tool calls across its public API and 500,000 sessions of its own coding agent, and published the results in February 2026 **VERIFIED C01**. Three findings sit uncomfortably next to the standard policy clause.

First, approval thins out as complexity rises. Human involvement is present in 87% of tool calls on minimal-complexity tasks — editing a line of code — and 67% on high-complexity tasks. The mechanism is structural rather than cultural: step-by-step approval becomes impractical as the number of steps grows. The control is weakest precisely where a mistake is most expensive. **VERIFIED C01**

Second, experience pushes operators away from per-action approval. New users of Claude Code run in full auto-approve mode in roughly 20% of sessions; by around 750 sessions of tenure that rises above 40%. Crucially, those same experienced users interrupt more often, not less — from about 5% of turns to about 9%. They are not abandoning oversight. They are swapping gate-keeping for monitoring. **VERIFIED C01**

Third, the lab that produced this data recommends against mandating the control. Anthropic's published guidance is explicit: oversight requirements that prescribe specific interaction patterns, such as requiring humans to approve every action, "will create friction without necessarily producing safety benefits," and the focus should be on whether humans are positioned to monitor and intervene **VERIFIED C01**. That is a model developer arguing against a rule that would be commercially convenient for it to accept.

Where oversight actually goes

Anthropic first-party measurement, public API and Claude Code. Source C01.

HUMAN INVOLVEMENT BY TASK COMPLEXITY

Minimal complexity — 87%

High complexity — 67%

20 points of oversight disappear as tasks get harder, not easier.

OPERATOR TENURE CHANGES THE OVERSIGHT STYLE

Auto-approve 20% new operators (<50 sessions)

Auto-approve >40%

experienced operators (~750 sessions)

5%

9%

interrupt rate per turn rises with tenure — monitoring replaces gate-keeping

Figure 1. Oversight does not vanish with experience; it changes shape. Policies that only recognize the gate-keeping form will record a control failure where an operator is in fact supervising well. Source C01.

Why this matters more in a regulated shop than in a software team

Anthropic's own caveat is the load-bearing one for regulated buyers: software is unusually amenable to supervisory oversight because outputs can be tested, compared and reviewed before release. In domains where verifying an agent's output requires the same expertise as producing it — adjudicating a claim, reading an image, approving a switching order — trust develops more slowly and the monitoring form of oversight is harder to build. **VERIFIED C01** Software engineering still accounts for nearly half of agentic tool calls on that API. Financial services, healthcare and cybersecurity appear, but as emerging rather than dominant usage.

The practical inversion. Stop asking "did a human approve this action?" Start asking four questions an examiner can actually test:

Four controls that replace the approval checkbox. Ariana Digital practitioner framing; evidence anchors in Source C01.		
CONTROL	THE QUESTION IT ANSWERS	EVIDENCE AN EXAMINER CAN INSPECT
Reversibility class	If this action is wrong, can it be undone, and by whom?	Every tool in the agent's registry tagged reversible, compensable or terminal. Terminal tools require a second factor, not a click.
Blast radius	How many records, dollars or customers can one agent run touch?	Per-run caps enforced at the tool layer, not in the prompt. Breach of cap halts the run and files a ticket.
Time-to-halt	How long between a bad signal and the agent stopping?	Measured in seconds, tested quarterly with an injected fault, logged like a disaster-recovery drill.
Evidence of intervention	Did a human actually change the outcome when it mattered?	Interrupt and agent-initiated stop rates by task class, trended. A flat line near zero is a finding, not a pass.

Ariana Digital take. We have never seen an agent program fail because a human forgot to click approve. We have seen them fail because nobody could say how far a single bad run could reach before somebody noticed. Write the blast radius into the tool definitions on day one; it is a two-week change before launch and a two-quarter change afterward.

3. Frontier ledger: what actually shipped, by lab

Equal editorial weight, unequal news volume. What follows is dated, first-party where available, and separates shipped product from scheduled or announced intent.

Frontier lab activity relevant to regulated enterprise buyers, as of Monday, September 21, 2026.			
LAB	SHIPPED AND DATED	GOVERNANCE SURFACE IT CREATES	SOURCE
Anthropic	Published field measurement of agent autonomy across 998,481 API tool calls and 500,000 coding sessions (February 18, 2026). Disclosed on September 17–18, 2026 that Claude leads roughly 26% of its internal R&D, up from none in February, with roughly 30,000 agents running at any time in August 2026.	Post-deployment monitoring metrics an enterprise can copy: reversibility share, safeguard coverage, agent-initiated stop rate. The R&D figure is a company disclosure about its own environment, not a benchmark.	VERIFIED C01 CITED C02
OpenAI	Workspace agents in ChatGPT, Codex-powered and generally available to Business, Enterprise, Edu and Teachers tenants (April 22, 2026). Expanded security partnership with CrowdStrike extending endpoint coverage to Codex agents (September 2, 2026). DevDay scheduled for September 29, 2026.	Admin-side scoping of tools and actions per user group, per-action approval for sensitive steps, a Compliance API exposing every agent's configuration, updates and runs, and the ability to suspend an agent. This is the closest thing on the market to an agent change-log a regulated auditor would recognize.	VERIFIED C04 CITED C21 CITED C22
Google DeepMind	Gemini Robotics 2, ER 2 and On-Device 2 (July 30, 2026), with published per-task success rates	The only frontier release this quarter that publishes its own failure rates task by task.	VERIFIED C03

LAB	SHIPPED AND DATED	GOVERNANCE SURFACE IT CREATES	SOURCE
	and the ASIMOV-Agentic safety benchmark. ER 2 available in Google AI Studio and in private preview on Gemini Enterprise Agent Platform.	ASIMOV-Agentic scores an agent's ability to refuse unsafe tool calls and to proactively request human intervention — a directly reusable acceptance criterion.	
SpaceXAI (xAI) and Cursor	Grok Bot for Enterprise (September 3, 2026) and Grok Bot for Procurement (September 4, 2026), with access, network and audit controls. Memory in Grok Build (September 16, 2026). Grok Voice Transcribe 2.0 (September 18, 2026). Grok Bot bundled into Cursor Pro and all Cursor Teams plans (August 26, 2026). Grok 4.6 distributed through Microsoft Foundry, Amazon Bedrock and Gemini Enterprise Agent Platform (August 2026).	Persistent always-on agents with their own compute, marketed by function — procurement first. Function-specific agents arrive with function-specific control requirements: segregation of duties, vendor master integrity, payment-file custody.	VERIFIED C05
NVIDIA (infrastructure layer)	Cosmos 3 world foundation model, GR00T N1.7 in early access with commercial licensing, and the Physical AI Data Factory Blueprint. FANUC and ABB integrating the Isaac platform. Filing shows supply commitments of roughly \$279 billion, up from roughly \$119 billion three months earlier	Synthetic training-data provenance becomes an auditable artifact. If a robot's behavior was learned in simulation, the simulation is now part of the validation record.	CITED C13, C23
	CITED C23		

4. Robotics reality check: read the failure rates, not the footage

The humanoid news cycle runs on video. The useful document this quarter is a bar chart. Google DeepMind published per-task success rates for Gemini Robotics 2 on July 30, 2026, using a single model checkpoint across three embodiments, and wrote plainly that multi-finger dexterous manipulation remains challenging

VERIFIED C03

Published success rates, Gemini Robotics 2 (July 30, 2026)

Single model checkpoint across three embodiments. Source C03.

GRIPPER — FRANKA DUO



WHOLE BODY — APOLLO 2 + INSPIRE HANDS



MULTI-FINGER — APOLLO 2 + SHARPAWAVE HANDS



Figure 2. The same model checkpoint scores 89.6% on precise insertion with a two-finger gripper and 32% on sweeping with a dustpan using a 22-degree-of-freedom hand. Constrained, jig-supported tasks are near production; open-world manipulation is not. Source C03.

The cause-and-effect an operations leader should take from this. Success rate is not a property of the robot. It is a property of the task's constraint. Precise insertion works because the fixture removes ambiguity. Picking from the floor drops to 45.7% because the floor has none. That implies the cheapest way to raise a robotics pilot's success rate is usually to change the environment, not the model — jigs, bins, fiducials, lighting, single-SKU lanes.

VERIFIED C03

What "deployed" currently means. Agility Robotics reports more than 65,000 operating hours for its Digit robot across nine customer facilities, naming GXO, Schaeffler, Toyota Motor Manufacturing Canada and Mercado Libre. BMW Group's AEON deployment in Germany is its first humanoid placement in Europe and targets high-voltage battery assembly by the end of 2026 — an announced target, not a delivered outcome. Reported deployments cluster in material handling, bin picking and inspection routes, not high-speed welding or stamping.

Reusable acceptance criterion. DeepMind released ASIMOV-Agentic alongside the models: a benchmark measuring whether the reasoning layer refuses unsafe tool calls from the action layer, predicts whether a task is possible, and proactively requests human intervention when uncertain.

VERIFIED C03

Any robotics or physical-AI vendor pitching a regulated site should be able to answer, in writing: what fraction of unsafe commands does the planner refuse, and how often does it ask for a human when it should?

Ariana Digital take. In a regulated plant, a 92% task success rate is not 92% good. It is an 8% exception queue that somebody has to staff, and the exception handler is usually the most experienced operator on the floor **VERIFIED C03**. Model the exception path and its labor cost in the business case before signing the pilot; that single line moves more pilots to a defensible yes-or-no than any model benchmark.

5. Financial services: the supervisor moved the goalposts, not the bar

WIN · CONSTRAINT · CONTROL

Win. Agentic deployment in banking has moved past proof-of-concept at the top of the market. JPMorgan Chase describes running more than 400 production AI use cases and has said it plans to deploy more capable agents during 2026 [CITED C16, company-reported](#). At the survey level, 40% of organizations with more than \$1 billion in annual revenue report scaling AI agents, up from 27% a year earlier, while smaller organizations were flat at 22% [CITED C10](#). The gap between large and small banks is widening, and it is a control-infrastructure gap more than a model-access gap.

Constraint. On April 17, 2026 the Federal Reserve, OCC and FDIC issued revised model risk management guidance — Fed SR 26-2 and OCC Bulletin 2026-13 — replacing SR 11-7 for the first time in roughly fifteen years. Generative and agentic AI were expressly placed *outside* the scope of that guidance, with the agencies stating that an institution's own risk management and governance practices should determine controls for systems not covered, and signaling a future request for information [CITED C06](#). That is not a green light. It means agentic systems are supervised under safety-and-soundness expectations and existing consumer-protection law without a purpose-built letter to point at during an exam.

Control that works this quarter. Build the agent control narrative your examiner cannot ask for yet. Map each agent to an existing, examinable regime rather than to "AI governance": third-party risk for the model and the marketplace, change management for prompts, skills and tool registries, records retention for agent transcripts, and complaint-handling for any customer-facing output. Then add the one artifact that is genuinely new — a tool registry with reversibility class and per-run blast-radius caps, exported monthly.

Scenario, 18 months out. When the agencies do issue AI-specific expectations, the most likely shape — given the direction of the April rewrite and the EU's own deferral to December 2027 [VERIFIED C15](#) — is outcome and monitoring oriented rather than prescriptive about interaction design. Institutions that spent this window building per-action approval queues will have compliance theater with a headcount attached. Institutions that spent it building monitoring, halt and rollback telemetry will have evidence. The asymmetry favors telemetry: it is useful whether or not the rule arrives.

Ariana Digital take. The highest-yield first agent in a mid-size bank is rarely customer-facing. It is dispute intake triage or KYC remediation packaging — high volume, fully reversible, already staffed by a queue with an existing quality-assurance sample. You get a measurable baseline, an existing control environment to plug into, and no new regulatory surface.

6. Healthcare: the FDA is asking the right question before the market answers it

WIN · CONSTRAINT · CONTROL

Win. Ambient documentation is the first genuinely scaled clinical AI deployment. Kaiser Permanente reports ambient AI clinical documentation running across roughly 40 hospitals and eight states with more than 10,000 clinicians, governed by a standing quality-assurance program rather than a pilot committee [CITED C17, operator-reported](#). The transferable lesson is not the tool. It is that the deployment shipped with a permanent QA function attached to it from the start.

Constraint. FDA's device center has published a discussion paper on generative-AI-enabled device software functions, describing a two-axis risk framework based on the degree and the independence of device activity, and naming confabulation, uncertain intended-use boundaries, limited visibility into third-party foundation models, and performance degradation across the product life cycle as distinct risks. The comment window on docket FDA-2026-N-7874 is scheduled to close on October 19, 2026 [CITED C07](#). Two axes — degree and independence of activity — is almost exactly the risk-and-autonomy plane the Anthropic field study plotted empirically [VERIFIED C01](#). Regulator and lab converged on the same two variables from opposite directions, which is a reasonable signal that those are the variables to instrument.

Control that works this quarter. Classify every clinical or clinical-adjacent agent on those two axes now, in writing, with a named owner per cell. Where independence is high and the activity touches diagnosis or treatment selection, assume a device conversation is coming. Where independence is high but the activity is administrative — prior authorization assembly, denial appeal packaging, coding suggestion — document the human decision point that stands between the agent and the patient record, and measure how often it actually changes the output.

The number nobody wants to publish. Health systems will quote turnaround-time improvements on appeals and revenue-cycle work all day. Very few publish the override rate: how often the reviewing clinician or coder materially changed what the agent produced. That single metric is the difference between an agent that works and an agent whose reviewer has stopped reading. Anthropic's field data shows why it matters — interrupt rate is what distinguishes active monitoring from rubber-stamping [VERIFIED C01](#). Track override rate by reviewer and by task class from week one.

Ariana Digital take. If your ambient or revenue-cycle agent has an override rate under about 2% and falling, do not celebrate. Run a seeded-error test: inject known-bad drafts into the review queue at a low rate and measure detection. If reviewers miss them, you do not have a human-in-the-loop control; you have a latency tax. This is a one-sprint diagnostic and it changes board conversations.

7. Manufacturing: the model is not the bottleneck, the data factory is

WIN · CONSTRAINT · CONTROL

Win. Constrained industrial tasks are where physical AI is actually paying. Reported ranges on production-ready data cluster at 15–30% downtime reduction from predictive maintenance and 40–60% inspection-labor reduction from quality triage, with Siemens reporting roughly a 20% throughput increase at its Erlangen electronics factory under an AI-driven adaptive manufacturing approach [CITED C19, company-reported](#). Note the qualifier in the source: *on ready data*. That phrase is doing most of the work.

Constraint. Two of them. First, dexterity: the published Gemini Robotics 2 figures show open-world manipulation still in the 32–46% band while jig-constrained insertion clears 89% [VERIFIED C03](#). Second, provenance: NVIDIA's Physical AI Data Factory Blueprint and Cosmos 3 world model make synthetic training data the default path to volume [CITED C13](#). Synthetic data is a genuine unlock for sample efficiency and a genuine problem for a plant that must explain, after an incident, why a machine behaved as it did. FANUC and ABB integrating the Isaac platform means this question arrives on ordinary industrial robots, not just humanoids.

Control that works this quarter. Treat the simulator as a controlled document. Version the scene, the domain-randomization ranges and the policy checkpoint together, and keep the triple immutable for the life of the deployed behavior. When a safety engineer asks what the robot was trained to expect, the answer should be a hash, not a conversation.

Risk and reward, stated plainly. The reward case for a constrained-task robotics cell is strong and getting stronger: the task is repetitive, the fixture is cheap, and the success rate is measurable per cycle. The risk case is almost entirely about the exception path and the integration labor. Reported deployments still require on-site vendor engineering and custom environment preparation. Price the pilot with the vendor's field engineer in it, because you will have one.

Ariana Digital take. Ask every physical-AI vendor for two numbers before the technical deep-dive: mean time between human interventions on your actual line, and the cost of the environment changes needed to reach their quoted success rate. Vendors who can answer both in one call are running real deployments. The rest are running demos.

8. Energy and utilities: AI is now a load-side reliability problem as much as an operations tool

WIN · CONSTRAINT · CONTROL

Win. Utilities have quietly produced some of the most durable AI operating results anywhere, because the use cases are forecast-shaped and the feedback loop is physical. National Grid reports a data-driven vegetation management program that cut tree-related events by roughly 30% and customer interruptions by roughly 38%; Duke Energy reports its self-healing grid program has prevented more than 1.5 million customer outages [CITED C18, operator-reported](#). These are optimization and prediction wins with human operators still holding the switching authority — which is exactly why they cleared reliability review.

Constraint. The sector's headline AI risk this year is not an agent misbehaving; it is AI's own demand behaving badly. On May 4, 2026 NERC issued a rare Level 3 "Essential Actions" alert after repeated events in which 1,000 MW or more of computational load dropped off the bulk power system within seconds — a magnitude comparable to a large generator trip — and set seven required actions for transmission planners, transmission operators, planning coordinators and balancing authorities, with an August 3, 2026 response deadline [VERIFIED C08](#). Separately, FERC issued six orders to regional grid operators on June 18, 2026 addressing large-load interconnection, generation adequacy, queue management and cost allocation; the 30-day and 60-day response windows had both elapsed by September 3, 2026 [CITED C09](#).

Control that works this quarter. If your organization is siting compute, your interconnection application is now a reliability commitment. Treat ramp-rate limits, ride-through behavior and curtailment response as engineering requirements owned by the same team that owns the AI roadmap — not as a facilities footnote. If your organization is a utility, the control is the inverse: model large computational load as a dispatchable, fast-moving entity in planning studies rather than as flat baseload.

Cause and effect that most AI strategies miss. An enterprise AI program with an on-premises or colocated inference footprint now has a regulatory dependency it did not have eighteen months ago. The same board deck that shows agent adoption curves should show the interconnection timeline, because the second one determines whether the first is executable. That dependency runs in both directions: utilities are simultaneously the constraint on AI buildout and among the most disciplined operational adopters of AI in the economy.

Ariana Digital take. For industrial and energy clients we now run the AI capacity plan and the electrical interconnection plan as a single workstream with one dependency map. It surfaces a conflict roughly two quarters earlier than running them separately, and two quarters is usually the difference between a re-scope and a write-off.

Ariana Digital reference shape. Each layer emits one auditable artifact.

1 · Intent and task contract What the agent may attempt, what counts as done, what it must escalate.	ARTIFACT Scoped task definition
2 · Tool registry with reversibility class Every tool tagged reversible, compensable or terminal. Per-run blast-radius caps.	ARTIFACT Signed tool inventory
3 · Execution plane Model, orchestration, memory. Interchangeable. Not where your control lives.	ARTIFACT Immutable run log
4 · Monitoring and halt plane Time-to-halt, interrupt rate, agent-initiated stops, injected-fault drills.	ARTIFACT Halt-drill record
5 · Evidence plane Mapped to third-party risk, change management, records retention, complaints.	ARTIFACT Monthly control pack

Applying the stack to one high-volume, reversible workflow. Control design is Ariana Digital practitioner work; the oversight metrics are drawn from Source C01 and the two-axis framing from Source C07.

LAYER	CONCRETE DECISION	FAILURE IT PREVENTS
Task contract	Agent may read a denial letter, assemble supporting documentation and draft an appeal. It may not submit, and it may not alter the clinical record.	Scope creep from drafting into acting, the single most common cause of an unplanned device or records conversation.
Tool registry	Document retrieval: reversible. Draft generation: reversible.	An agent that can file a claim it should not have filed. Removing

LAYER	CONCRETE DECISION	FAILURE IT PREVENTS
	Submission API: terminal, removed from the agent's registry entirely.	the tool beats approving its use.
Execution plane	Any competent frontier model with enterprise admin controls. Vendor-neutral by design; re-evaluated quarterly on cost and quality.	Architectural lock-in to a single lab at the exact moment distribution is consolidating onto marketplaces.
Monitoring and halt	Override rate tracked per reviewer and per denial type. Seeded-error injection at a low rate. Halt on cap breach within seconds.	Reviewer fatigue silently converting a control into a rubber stamp.
Evidence	Monthly pack: run volume, override rate, seeded-error detection rate, halt drills, tool-registry diffs.	Arriving at an audit with model documentation and no operating evidence.

Where the market scoreboard actually stands

Deployment is not adoption, and adoption is not value

Sources C10 (McKinsey), C11 (Deloitte), C12 (Stanford AI Index, as reported).

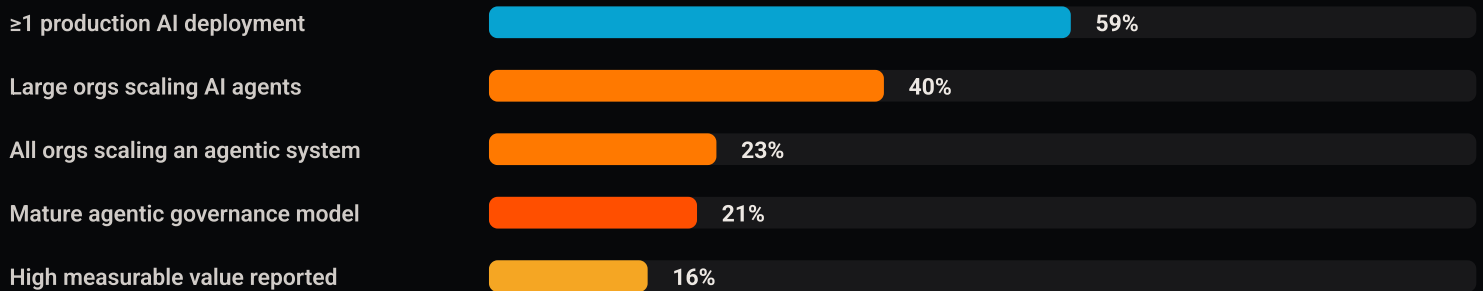


Figure 4. The 59-to-16 gap is the market. Note that governance maturity (21%) and measurable value (16%) sit almost on top of each other — organizations that can govern an agent are broadly the ones that can measure it, because both require the same instrumentation. Sources C10, C11, C12.

10. Scenario planning and risk-reward

Three plausible paths through the next four quarters, with the leading indicator to watch for each. Scenarios are Ariana Digital analysis, anchored to the cited regulatory and market evidence.

SCENARIO	WHAT IT LOOKS LIKE	LEADING INDICATOR	POSITION TO TAKE NOW
Quiet consolidation most likely	Agent capability keeps improving; regulators keep deferring specifics. The EU's December 2027 deadline and the US banking RFI both land as monitoring-oriented expectations VERIFIED C15 CITED C06 Winners are decided by integration quality, not model choice.	Marketplace distribution share. When a lab's enterprise revenue arrives mostly through hyperscaler marketplaces, contractual controls standardize and differentiation moves to the workflow.	Build vendor-neutral at layer 3. Invest in layers 2, 4 and 5, which survive any model swap.
Incident-driven tightening	A publicized agent failure in a regulated workflow — most plausibly a payments, claims or clinical-documentation error at scale — triggers prescriptive rules written quickly and applied broadly.	The first enforcement action or supervisory finding that names an agent rather than a model. Watch complaint volumes on agent-touched channels.	Be able to produce a halt-drill record and a tool registry diff within one business day. That capability is the whole defense.

SCENARIO	WHAT IT LOOKS LIKE	LEADING INDICATOR	POSITION TO TAKE NOW
Value stall	Adoption continues while measurable value stays near the 16% band CITED C12 budgets rotate to fewer, deeper deployments. Forecasts of large-scale agentic project cancellation by the end of 2027 CITED C20, forecast partially materialize.	Internal: ratio of agents in production to agents with a named business owner and a baseline metric. Above roughly three-to-one, a cull is coming.	Kill the unowned agents yourself, early and visibly. Concentrate spend on the two workflows with existing quality-assurance baselines.

II. Practitioner FAQ and did-you-know

Our policy requires human approval of every agent action. Are you saying we should remove it?

No — we are saying do not let it be your only control, and do not assume it is being exercised. Keep per-action approval where the action is terminal and low-volume. Everywhere else, measure whether approval is actually changing outcomes. The field evidence is that approval density falls as complexity rises **VERIFIED C01** so a policy that assumes uniform approval is describing a system you do not have.

What is the single cheapest control to add this month?

Per-run blast-radius caps enforced at the tool layer. Not in the system prompt — at the tool layer, where a model cannot argue with them. Most teams can ship this in one sprint, and it converts an unbounded failure into a bounded one.

How do we evaluate a robotics or physical-AI vendor without a robotics team?

Ask for three numbers on your task, not theirs: per-cycle success rate, mean time between human interventions, and the environment changes required to reach the quoted rate. Then ask what fraction of unsafe commands the planning layer refuses — the ASIMOV-Agentic benchmark released with Gemini Robotics 2 gives that question a public reference point **VERIFIED C03**.

Did you know? Agents stop themselves more often than people stop them.

On the most complex tasks, Claude Code asked for clarification more than twice as often as humans interrupted it, and the top reason it stopped was to present the user with a choice between proposed approaches **VERIFIED C01**. Agent-initiated stops are a real oversight channel, and almost no enterprise control framework counts them. If your monitoring does not distinguish an agent that paused from an agent that failed, you are discarding your best early-warning signal.

Did you know? Only about 0.8% of sampled agent actions appear irreversible.

That figure **VERIFIED C01** reframes the whole control problem. The governance question is not "how do we supervise everything?" but "have we correctly identified the 1% that cannot be undone, and is it behind a different kind of gate?" Most organizations have never produced that list.

Which frontier vendor should a regulated mid-market firm standardize on?

On current evidence, none — standardize the control plane, not the model. All four major labs now ship enterprise administrative controls of broadly comparable shape: scoped tool access, per-action approval for sensitive steps, audit surfaces and suspension **VERIFIED C04, C05**. Model quality moves quarterly; your tool registry, halt plane and evidence pack should not have to.

Where Ariana Digital helps

AEGIS Diagnostic — a two-week read on where your agents actually stand against the four controls in Section 2, with the tool registry and blast-radius map produced as deliverables.

AEGIS Build — implementation of the intervention-capable stack in Figure 3 on one regulated workflow, instrumented end to end.

AEGIS Run — the monthly evidence pack and halt-drill cadence your examiner, auditor or board will ask for. AEGIS is the Agentic Enterprise Governance and Intelligence Standard.

[Get the AI Readiness Brief](#) · [Read the governance overview](#) · [Book a Diagnostic](#)

Ariana.Digital

Design, data, and emerging technology, from strategy through execution. Talent building and exceptional omni-channel customer experiences for companies who like to win.

SERVICES

Diagnose

Build

Run

Verticals

PRODUCTS

myndQ for Business

myndQ TalentHub

myndQ.ai

AI Success Pack

COMPANY

Industry Journeys

Reference Architecture Studio

About

AI Governance

All insights

2026 AI Readiness Brief

Contact

Contractor bench

FOLLOW

LinkedIn (Ariana.Digital)

LinkedIn (myndQ)

YouTube

