

RESEARCH REPORT

THE TRUST DIVIDEND

Why building trust in one AI agent can pay off in the next



Atlantic Re:think

×

salesforce

CONTENTS

Introduction: The New Enterprise Constraint	2
1. Defining What Agents Are Allowed to Do	4
2. Making Trust Part of the Foundation	8
3. Fixing the Bottlenecks That Slow AI Agents Down	11
4. Turning Trust Into a Competitive Advantage	15
5. Focusing Human Expertise Where It Matters Most	19
Conclusion: The Trustworthy Enterprise Moves Faster	23
Methodology	24

INTRODUCTION

THE NEW ENTERPRISE CONSTRAINT

The first AI-agent deployment tests whether an organization can trust a system to act. The next tests whether it can carry that trust forward.

You can't trust an agent to always be right—it's probabilistic, so the same setup can act differently each time. What you can build is the system around it: knowing how far a wrong action can reach, catching an agent that steps outside its limits, defining when a human takes over, and naming who's accountable when it fails. Build that system once and it becomes reusable infrastructure—helping you move faster and with more confidence. That's the **Trust Dividend**.

Evidence from testing, access controls, escalation paths, and clear accountability can give a comparable deployment a foundation to build on. Each new use case still needs scrutiny. The opportunity is to carry useful work forward and focus fresh review on what has changed.

That requires a practical definition of trust: what an agent may do, the evidence supporting that authority, the conditions for human intervention, and who is accountable for the outcome. As agents act across enterprise systems, these decisions shape both the risks organizations take and the work required to expand.

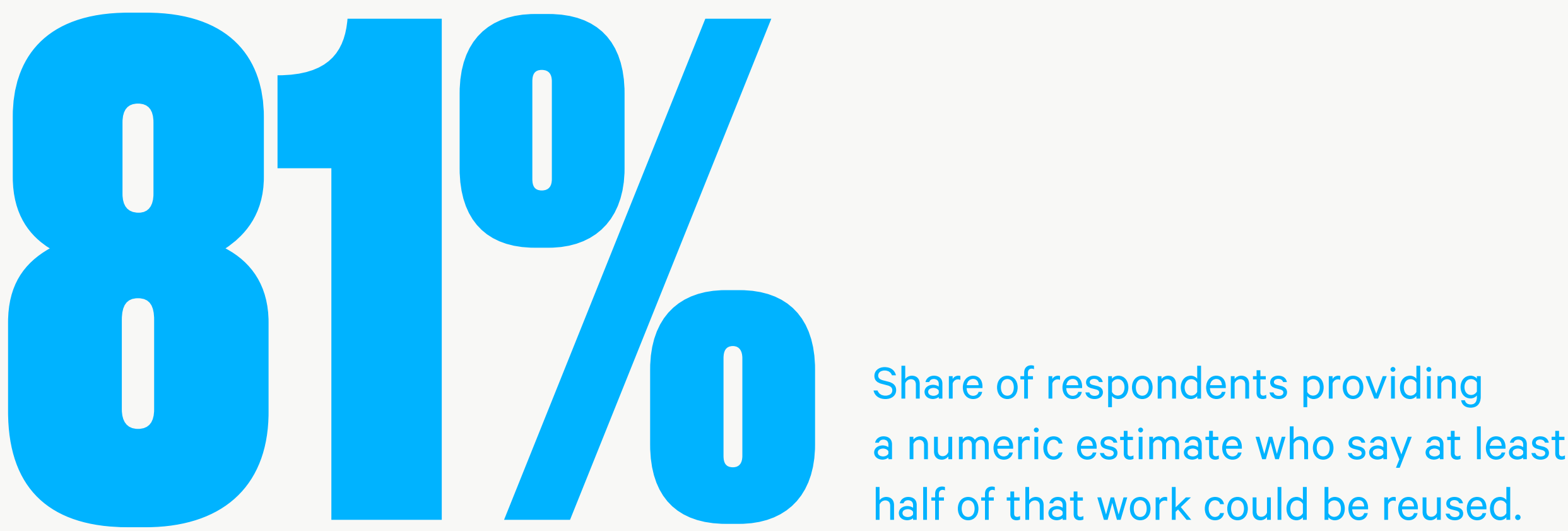
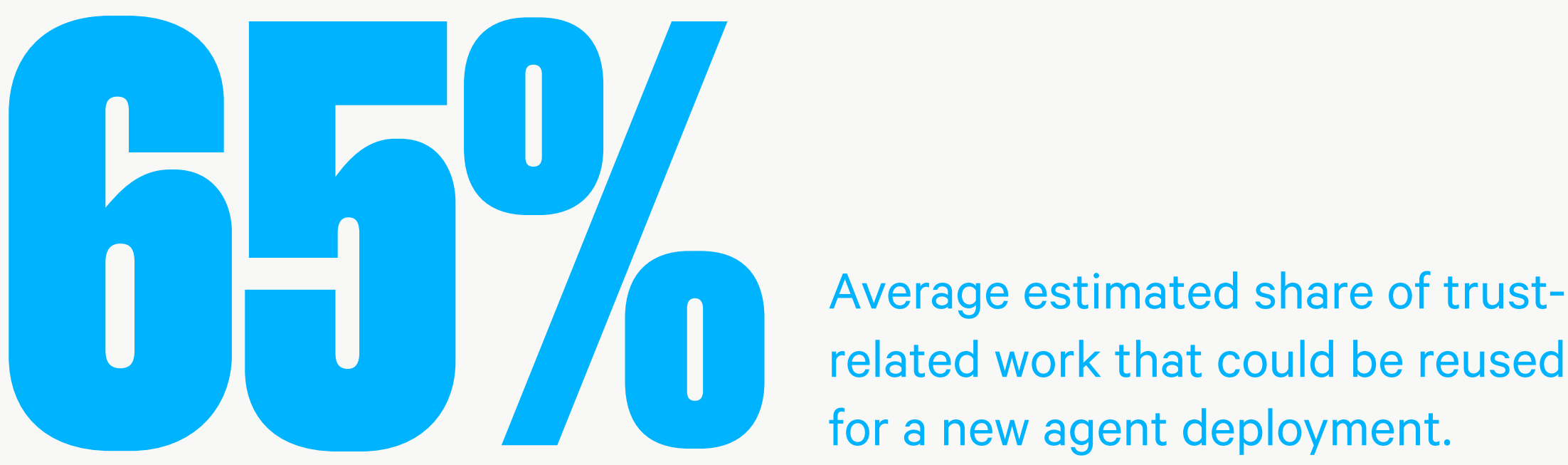
Atlantic Insights and Salesforce surveyed 300 U.S. business and IT decision-makers at organizations with at least 300 employees. All are director-level or above and report an AI agent in production. A diary study and interviews with Salesforce executives add accounts of how these decisions work in practice.

Reuse already has a substantial place in respondents' plans. Among those able to estimate it, an average of 65 percent of trust-related work could be reused for a new agent deployment; 81 percent say at least half could carry forward. These are estimates of reusable work, not measured reductions in deployment time.

Confidence and constraint coexist. Sixty-four percent of respondents giving a numeric answer express high confidence in expanding agent autonomy without materially increasing risk. Yet 62 percent of the full sample often or very often encounter demand for more agent authority than current trust capabilities can support.

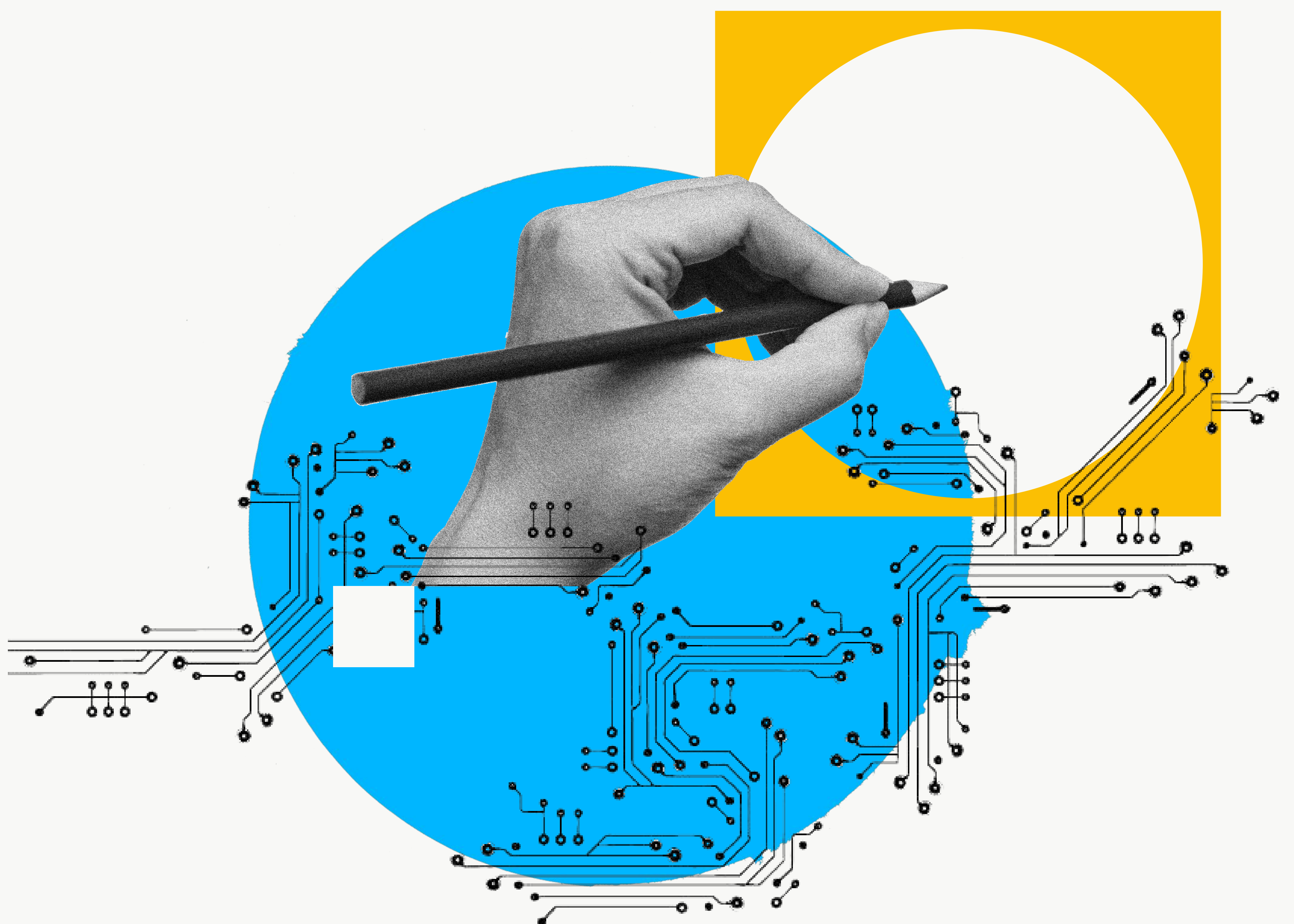
Trust requirements can also help work move. Seventy-one percent say those requirements speed deployment of another agent using a similar technical or operating pattern. But 77 percent report that a requirement surfacing later than expected has at least sometimes prompted a pause, rework, or additional review. Together, these findings make predictability and reuse practical priorities for leaders.

Across the next five chapters, we trace how trust moves from a one-time gate to a compounding advantage—starting with the permissions and infrastructure that make an agent's authority defensible, through the real bottlenecks that emerge as agents scale across platforms. From there, we show how that same trust work becomes a source of speed and competitive edge, and where human judgment still belongs when the stakes are highest. Together, these chapters build the case for the Trust Dividend: the return organizations earn when trust becomes reusable infrastructure rather than a recurring cost.



CHAPTER 1

DEFINING WHAT AGENTS ARE ALLOWED TO DO



The defining question of the agentic enterprise is permission. What's an agent allowed to do? With what data? Under what conditions? With what fallback path if the system is uncertain or the stakes are high?

For a CRO, this may be the difference between an agent drafting a sales follow-up and advancing a deal stage. For a marketer, it may be the difference between generating campaign variants and launching personalized outreach. For a CTO or CISO, it may be the difference between read access and write access across systems of record.

Sixty-four percent of survey respondents rate their confidence between an 8 to 10 (on a 0-to-10 scale, where 10 represents the highest level of confidence) on whether their organization can expand agent autonomy over the next 12 months without materially increasing risk. But the same data set complicates that confidence: 62 percent say their organization often or very often encounters situations where business leaders want an agent to have more scope or autonomy than current trust capabilities can support.

Experience appears to matter, but not as a simple cure. Organizations with 12 or more months of production-agent experience report higher confidence that they can expand agent autonomy without materially increasing risk: a mean confidence score of 8.2 versus 7.6 among organizations with less than 12 months of experience.

How confident are you that your organization can expand the autonomy of its AI agents over the next 12 months without materially increasing organizational or customer risk?

Among organizations with 12 or more months of production-agent experience:

8.2

“I think when you cross the threshold of competency to the point where it’s at or above human performance, that’s the first time where you go, ‘Now I have to confront the decision of, do I let it be autonomous?’”

JOE INZERILLO, PRESIDENT, ENTERPRISE & AI TECHNOLOGY, SALESFORCE

Inzerillo’s point is not that autonomy turns on all at once. In the same exchange, he describes mature adopters as earning the right to go fast by ramping autonomy gradually—starting with small shares of traffic, watching performance, and expanding only when the evidence supports it.

That tension—high confidence, persistent limits—suggests that the next wave of enterprise AI competition will center the operating discipline required to move from confidence to delegated action.

Evaluated on a 0-10 point scale where 10 is the highest level of confidence

Among organizations with less than 12 months of production-agent experience:

vs.

7.6

A DAY IN THE LIFE

OF AN AGENT OPERATOR

The Banking Agent That Can Advise, But Not Act

A systems administrator in a regulated banking environment described a reference agent that helps employees find answers from internal datasets. The agent can retrieve and synthesize information quickly, but its authority stops short of customer-facing communications or system-changing actions.

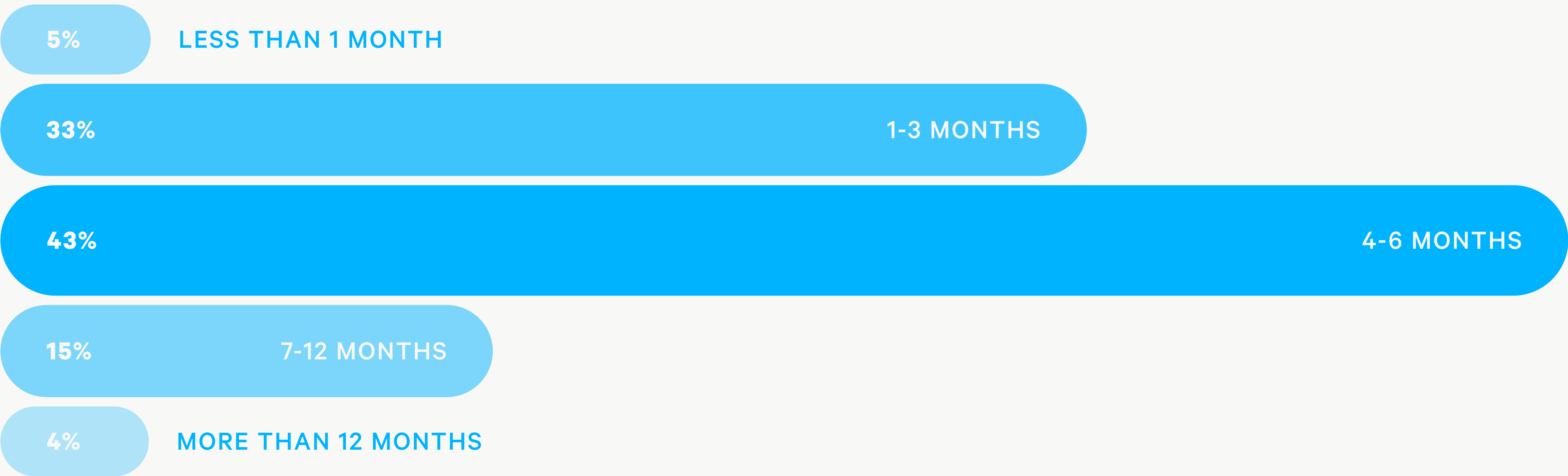
The participant's point was practical: the human was not a backup reviewer called in after the system acted. The human was the authority that

kept regulated decisions from crossing into customer-facing action. One workflow, the human did not re-enter after the fact; the human was already there as the decision-maker.

Moral of the story: Agent autonomy is not a single setting. Organizations define it by domain, consequence, reversibility, and regulatory exposure. In this case, the agent is trusted as an advisor to employees with no customer-facing capabilities.

This is not a generic adoption story. It’s a delegation story. Agents are being authorized to act and given greater operating authority. By operating authority, we mean what an agent can access, what actions it can take, where those actions occur, and whether the work expands into new functions or higher-consequence settings. But authorization is slow and uneven.

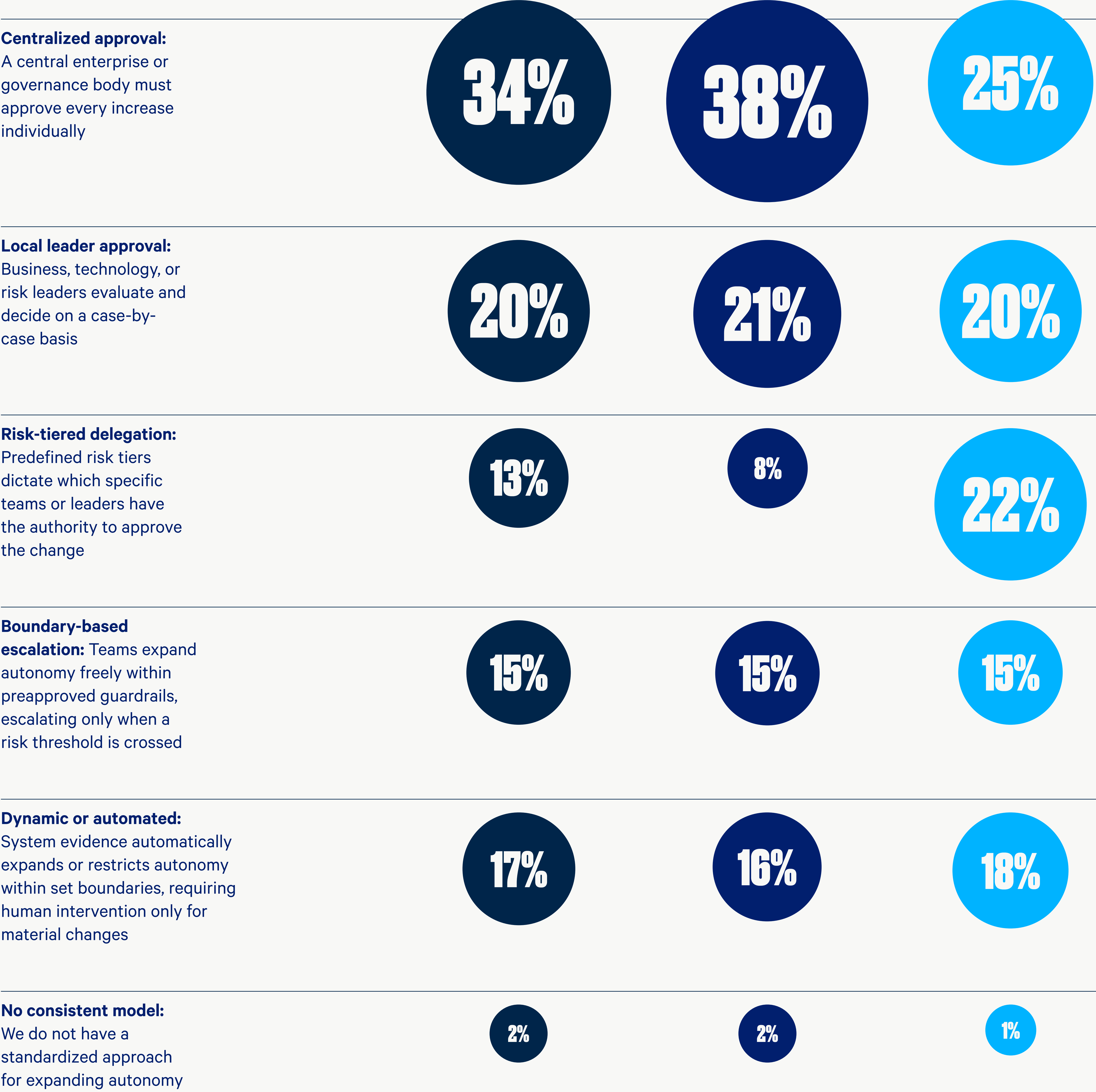
Approximately how much time passed between an AI agent satisfying the relevant trust requirements and receiving greater operating authority?



Base: Respondents whose most recently deployed agent received greater authority

Findings also show how organizations decide whether to increase an agent's autonomy. The difference among job levels points to a useful executive tension: **leaders at the top may see the enterprise risk of autonomy more clearly, while next-generation leaders may be closer to the operating models that could make autonomy more repeatable.**

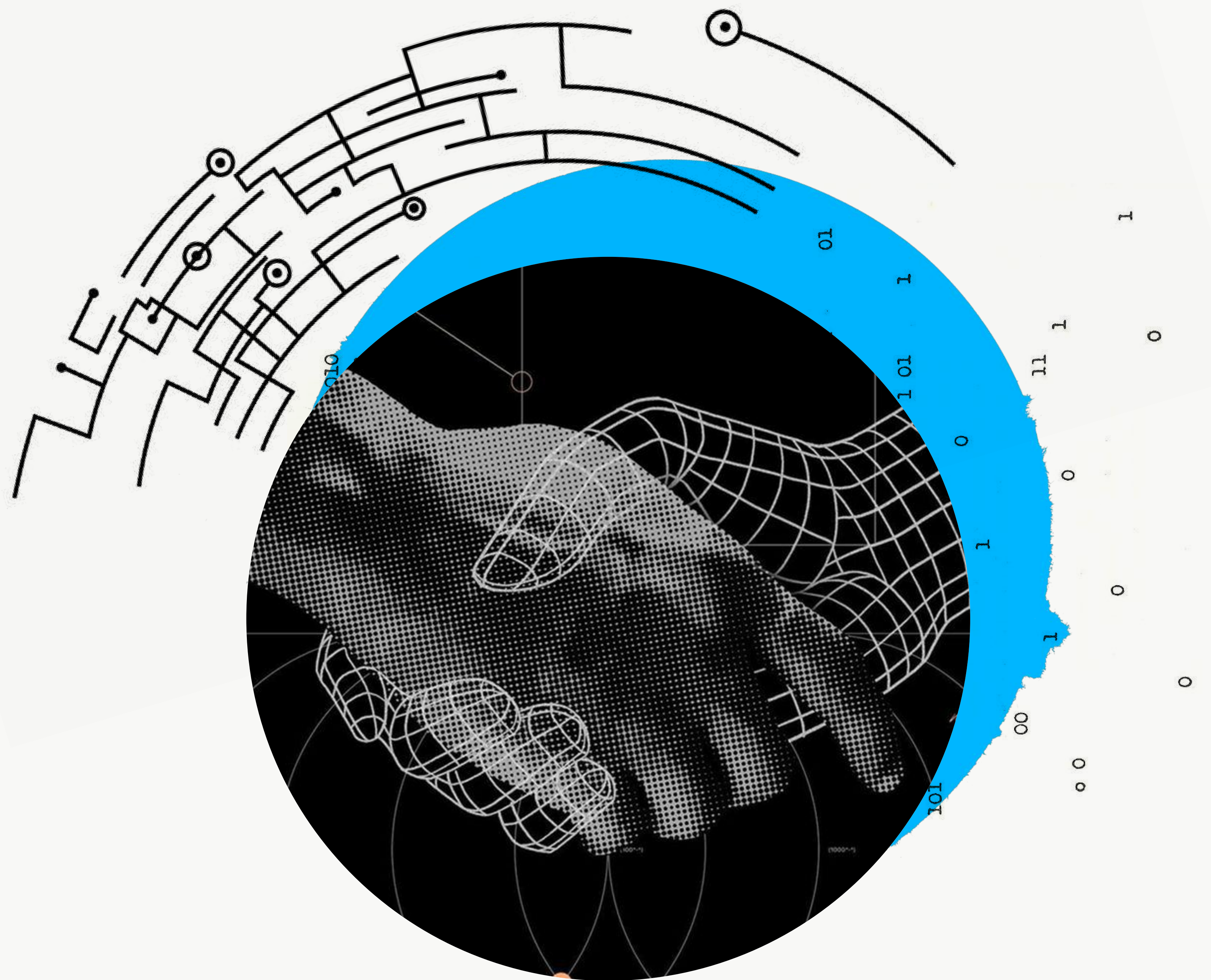
How is a decision to increase AI agent autonomy typically made?



The operating opportunity is to make the basis for an autonomy decision explicit: the evidence required, the limits approved, and the conditions that trigger escalation. Recording those decisions gives the next team something to assess and adapt. A comparable workflow can build on prior testing and established boundaries, while changes in data, actions, or consequences receive fresh review.

CHAPTER 2

MAKING TRUST PART OF THE FOUNDATION



Trust infrastructure includes the data an agent is allowed to use, the identity under which it acts, the permissions it can inherit or request, the provenance of the evidence behind its recommendation, the runtime policies that constrain it, the monitoring that detects drift or anomalies, the escalation routes that bring in humans, and the accountability model that assigns responsibility after an action is taken. What makes that foundation reusable is a record of where its controls have been tested and which conditions must still hold.

A DAY IN THE LIFE

The Security Agent With a Do-Not-Touch List

A security engineer described an agentic workflow that can proactively block malicious indicators in customer environments. That level of autonomy is possible only because the team maintains an exclusion list: a set of critical systems and artifacts the agent is not allowed to touch.

In a recent incident, the exclusion list prevented the workflow from blocking critical infrastructure. Without that safeguard, the

OF AN AGENT OPERATOR

participant said, “We would have blocked critical infrastructure in a customer’s environment.”

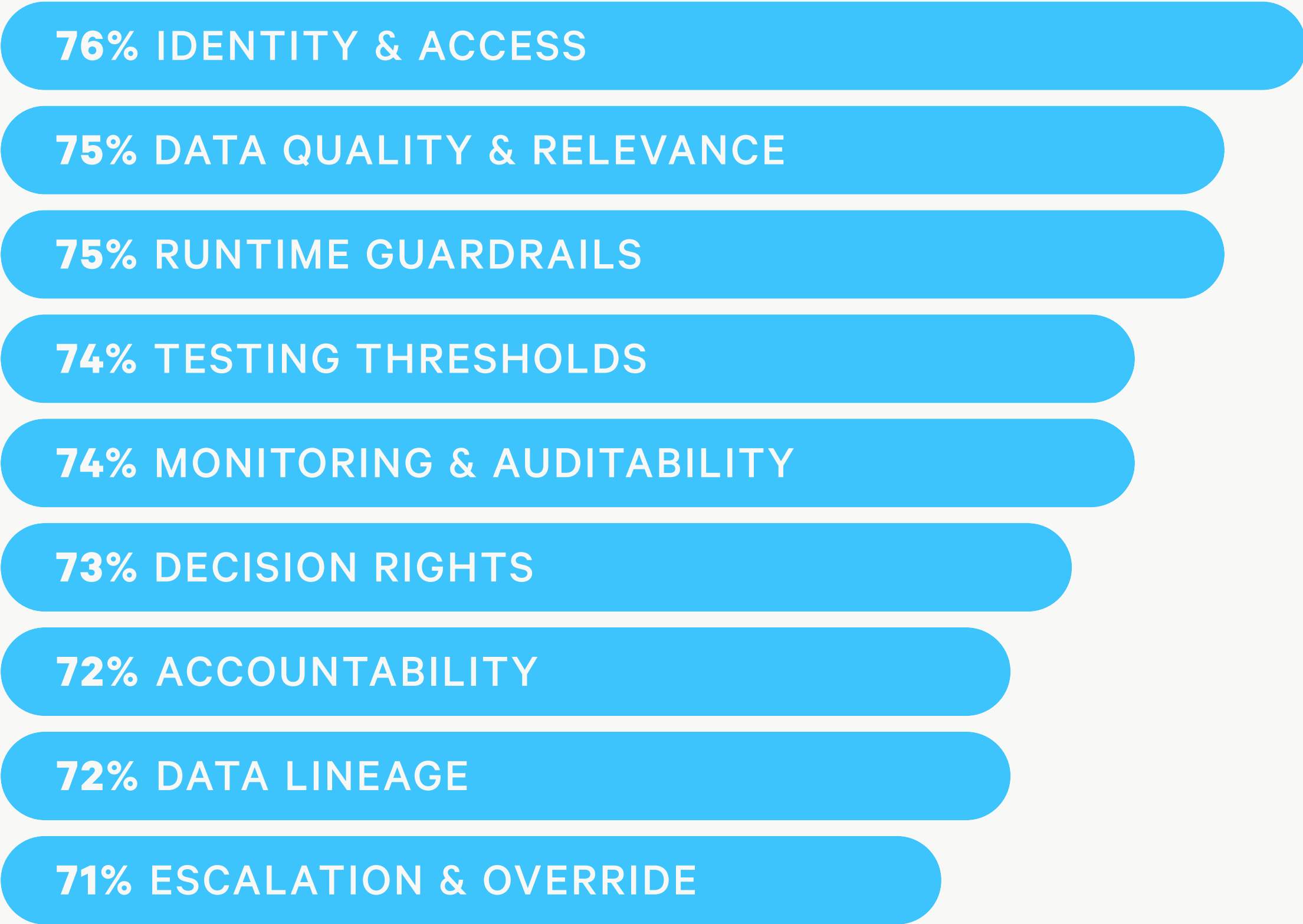


Moral of the story: Guardrails can be misunderstood as limits on AI. In production, strategic guardrails can be what make autonomy possible. The exclusion list does not slow the security agent down; it gives the organization enough confidence to let the agent act quickly inside defined boundaries.

The results of our survey proves that leaders aren’t starting from zero (trust). Across nine trust capabilities, roughly seven in ten respondents say their organization’s trust capabilities are either “standardized” or “platformized.” In the survey, “standardized” means formal standards are mandatory and consistently applied across deployments; “platformized” means core trust controls are built into centralized architecture.

How mature is your organization in each of the following aspects of its trust capabilities?

% share saying each trust capability is standardized or platformized



But deeper analysis illuminates the gap between standardization and platformization. Only about one-third of respondents report each capability is platformized. Identity and access leads at 37 percent; data quality sits at 36 percent; monitoring and auditability and escalation and override are each at 33 percent; runtime guardrails, accountability, and data lineage are each at 31 percent; and testing thresholds and decision rights are each at 30 percent. That’s the gap the Trust Dividend depends on: many enterprises have rules, but fewer have embedded those rules in systems that can travel across use cases.

Agent tenure makes that gap more revealing. Across all six controls — data quality, identity and access, data lineage, monitoring, decision rights, and accountability — organizations with 12 or more months of production-agent experience are consistently more likely to say the control is platformized, by margins of 11 to 19 percentage points. The longer an organization has lived with agents in production, the more trust appears to migrate from one-off policy into systematized architecture.

“Ethics is not just one person's job. The whole company should be able to raise questions, spot issues, and find ways to make the product better.”

PAULA GOLDMAN, EVP AND CHIEF ETHICAL AND HUMANE USE OFFICER, SALESFORCE

Timing is another test of whether trust functions as infrastructure or cleanup. Thirty-eight percent of respondents say trust requirements are upfront and predefined before any use-case approval. Twenty-one percent define requirements at use-case approval, combining enterprise baselines with tailored requirements. Together, 40 percent define trust requirements late or inconsistently—varying by function or use case, during pilot testing, or reactively after issues surface. That’s the cleanup problem: trust arrives when a team is already trying to move.

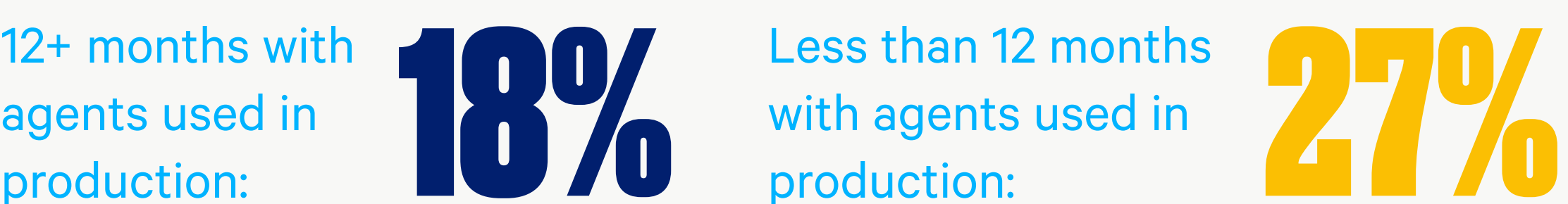
Here again, tenure matters. Experience teaches organizations that late trust is expensive trust.

When are trust requirements typically defined for AI agents entering production or gaining greater autonomy?

Upfront & predefined: Predefined enterprise standards and autonomy thresholds are established before any use-case approval



Variable: Timing varies substantially by function or use case



Findings around reuse put a number on that shift. Our respondents estimate that 65 percent of trust-related work can be reused when beginning a new AI-agent deployment. Eighty-one percent say at least half can be reused, while one-third say 75 percent or more can be reused. When an established deployment approach is extended to a comparable new use case, business unit, or geography, 60 percent say most or almost all of the prior trust evidence and approval can carry over. At the other end, 33 percent say about half or less carries over.

Which changes to your organization's trust practices would help the most to scale AI agents safely and quickly?

(ranked in order from Most Important to Least Important)

- 1 Improve the quality, accessibility, and governance of the data AI agents depend on
- 2 Use evidence from agents already in production more systematically to adjust permissions, controls, and autonomy
- 3 Define clearer thresholds for when an AI agent can act autonomously versus when human involvement is required
- 4 Make trust controls and evidence more reusable across AI-agent deployments rather than recreating them for each use case
- 5 Create more consistent enterprise-wide standards for testing, approving, and governing AI agents
- 6 Improve coordination across business, technology, risk, legal, and other teams involved in AI-agent decisions
- 7 Clarify who has authority to approve, expand, restrict, or pause an AI agent’s autonomy
- 8 Automate more trust controls, monitoring, and policy enforcement rather than relying on manual review
- 9 Establish clearer ownership and accountability for the business and customer outcomes AI agents produce
- 10 Improve the organization’s ability to measure whether trust capabilities are producing better business and customer outcomes

Reuse is the clearest measurable expression of the Trust Dividend. It turns trust from a project cost into an organizational asset. A permission model defined for one class of customer-service action can inform the next. A testing threshold used for one workflow can become a reference point for another. A data-lineage and auditability pattern can reduce the amount of proof required each time a team wants to extend an agent into a comparable setting.

CHAPTER 3

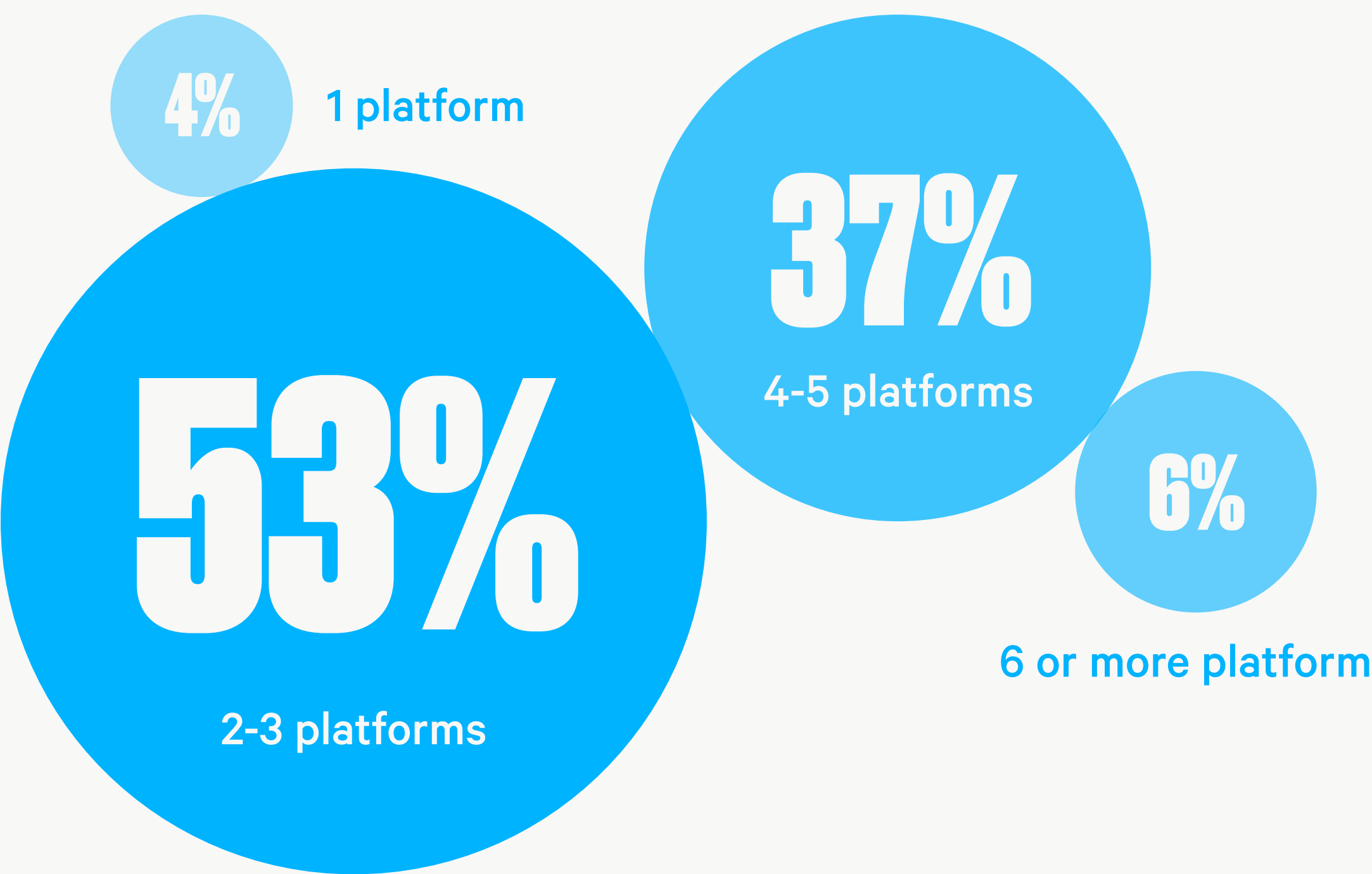
FIXING THE BOTTLENECKS THAT SLOW AI AGENTS DOWN



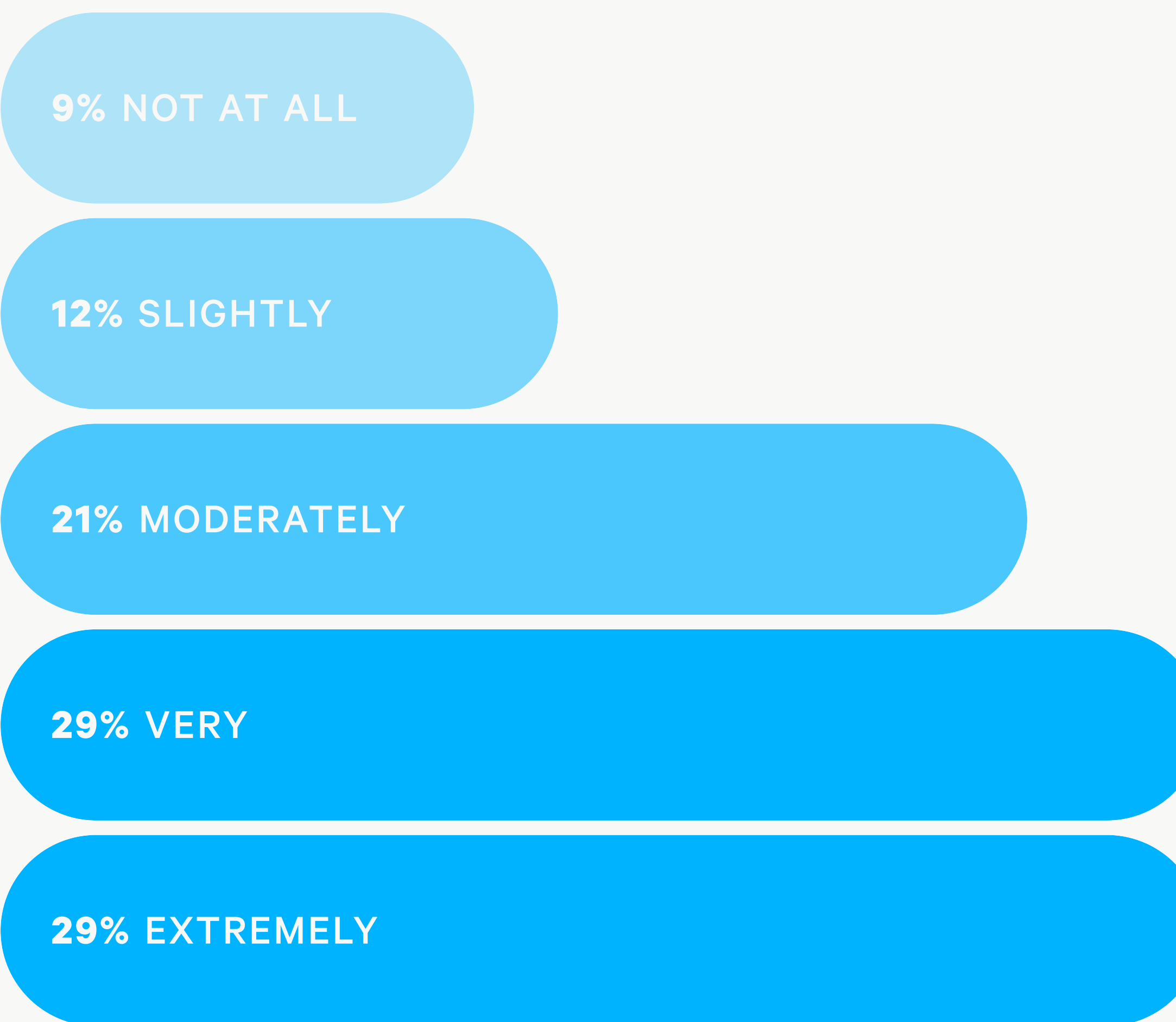
Reusable trust work does not guarantee a lighter operating burden. Teams still need to establish where existing controls apply, identify new requirements, and coordinate decisions across systems.

Ninety-six percent of respondents report using more than one platform or vendor environment for agents. Fifty-eight percent of all respondents say operating across multiple platforms makes consistent trust, governance, or security standards very or extremely difficult to maintain.

Approximately how many distinct platforms or vendor environments does your organization currently use to build, deploy, or run AI agents?



To what extent does having AI agents across multiple platforms or vendor environments make it harder to maintain consistent trust, governance, or security standards?



“It’s about data. It’s about the identities. It’s about permissions. Because in the past it was humans doing it, now it’s agents doing it and agents are relentless.”

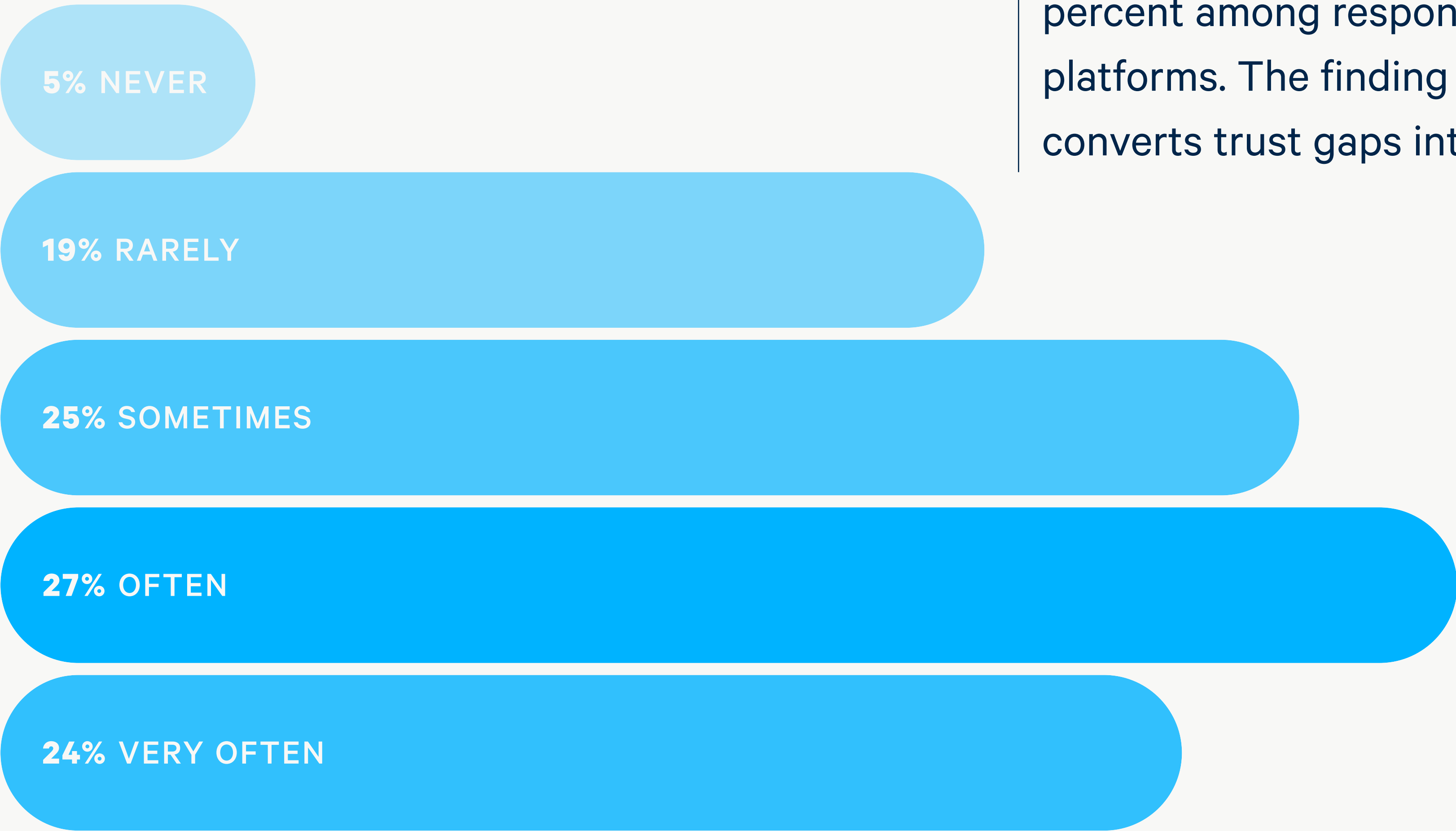
GAGAN GULATI, GM OF SECURITY, SALESFORCE

This is where security becomes central to the autonomy story. OWASP's 2025 Top 10 for LLM Applications identifies “excessive agency” as a major application risk: LLM-based systems can cause harm when they’re given too much functionality, too much permission, or too much autonomy relative to their intended role. That’s not a theoretical governance issue. It’s the technical condition under which agents move from answering to acting.

Gagan Gulati, Salesforce GM of Security, argues that an agent pursuing a goal may traverse systems in ways a traditional access model wasn’t built to govern. A human with too much access may still be constrained by attention, habit, and time. An agent can chain actions and search for another path. That means identity, permissions, lifecycle management, revocation, monitoring, and least privilege need to be designed for non-human actors.

Our survey results depict the operational cost when trust isn’t built into the work early enough. Seventy-seven percent of respondents say that in the past 12 months, an AI-agent deployment or planned expansion of autonomy has at least sometimes been paused, reworked, or sent through additional review because a trust requirement surfaced later than expected.

During the past 12 months, how often has an AI agent deployment or planned expansion of autonomy had to be paused, reworked or sent through additional review because a trust requirement surfaced later than expected?



Platform sprawl turns late trust requirements into operating drag. Among respondents using four or more agent platforms, 60 percent say late-surfacing requirements often or very often lead to pauses, rework, or additional review, compared with 45 percent among respondents using one to three platforms. The finding shows how complexity converts trust gaps into delay.

A DAY IN THE LIFE

OF AN AGENT OPERATOR

The Business Case for a Narrower Launch

A product manager at a software company described a reconciliation agent whose initial design allowed small variances to pass while flagging larger anomalies. Extending it to trust accounting exposed a gap between what the agent can do and what the workflow requires.

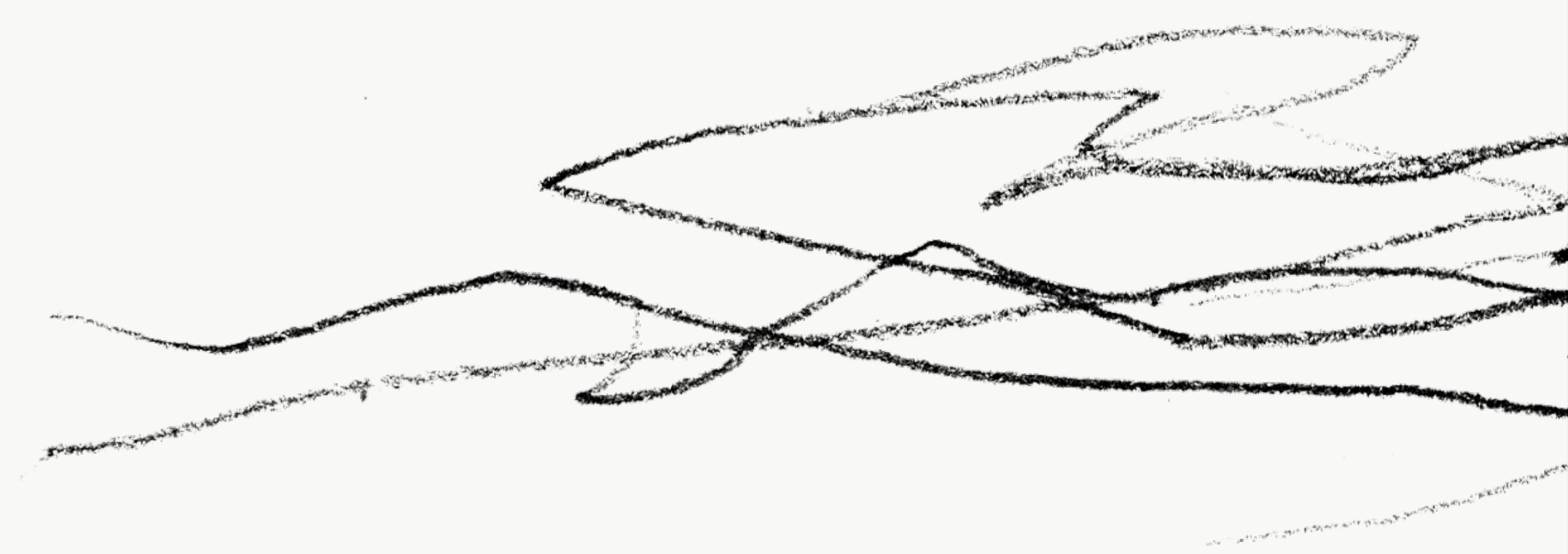
For that use case, the participant explained, every client ledger needs to balance to the penny, with discrepancies corrected and explained before proceeding. The team determined that the agent couldn't meet that

standard. So it limited the agent to operating-account reconciliation and shelved the trust-accounting application—a core workflow for its users. Adding safeguards took about two weeks of engineering work.

Moral of the story: The consequence is a narrower product and more work before its scope can expand. The case shows why requirements need to be specific to the workflow: evidence that supports automation in one setting may be insufficient in the next.

This friction isn’t limited to isolated projects. Compared with their earliest production agent deployments, 72 percent of all respondents say the overall effort required to satisfy trust requirements for a comparable new deployment has increased. At the same time, 71 percent of all respondents say ongoing human review per agent or workflow has increased as agent authority has grown. Only 15 percent say human review time has decreased.

Human review and resources are the hidden costs of agent sprawl.



As your organization increased the authority of its AI agents, how has the amount of ongoing human review required per agent or workflow been affected?

Increased substantially

1-3 agent platforms:

24%

4+ agent platforms:

52%

Compared with your organization’s earliest production AI agent deployments, how has the overall effort required to satisfy trust requirements for a comparable new deployment changed?

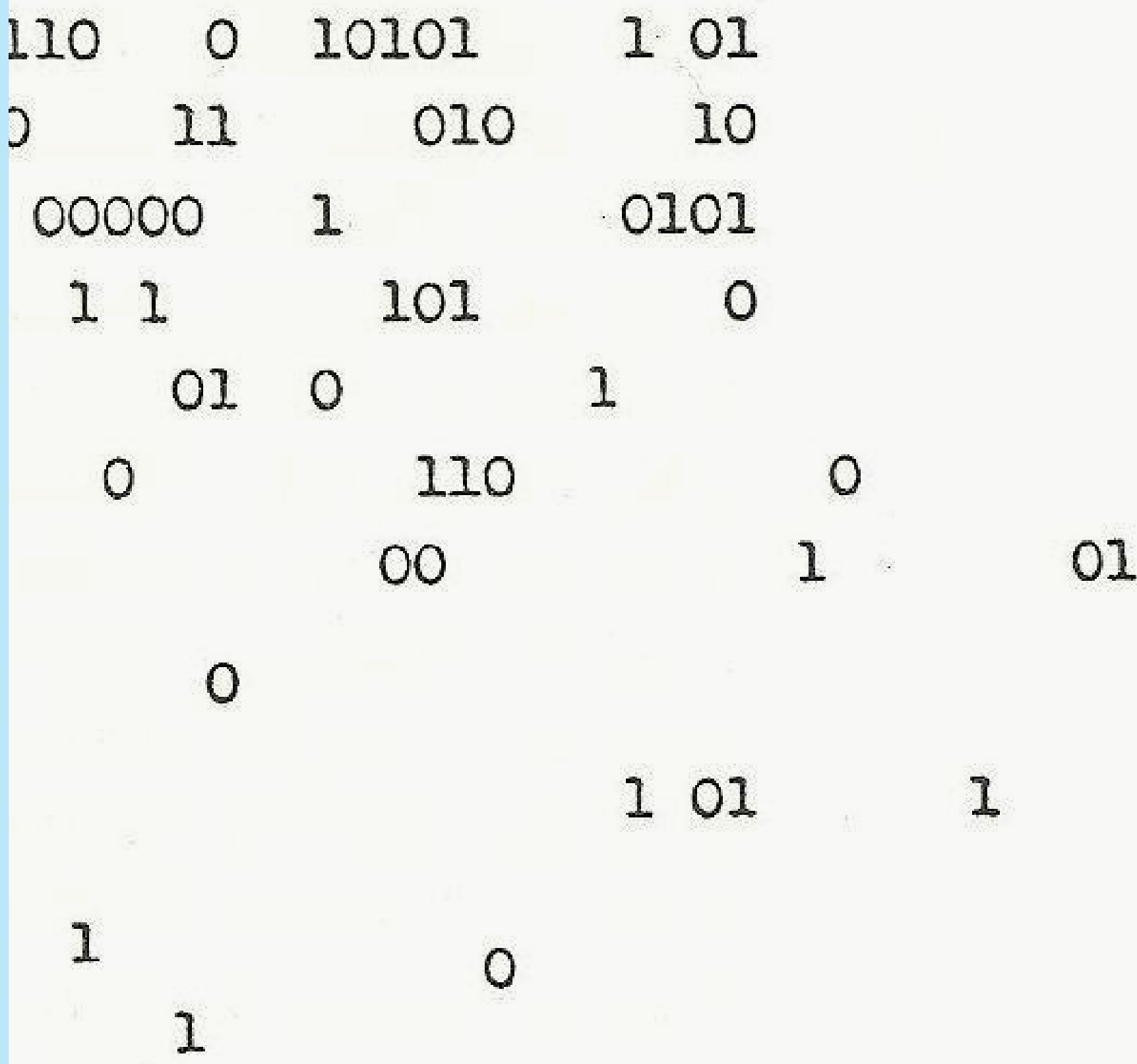
Increased substantially

1-3 agent platforms:

28%

4+ agent platforms:

43%

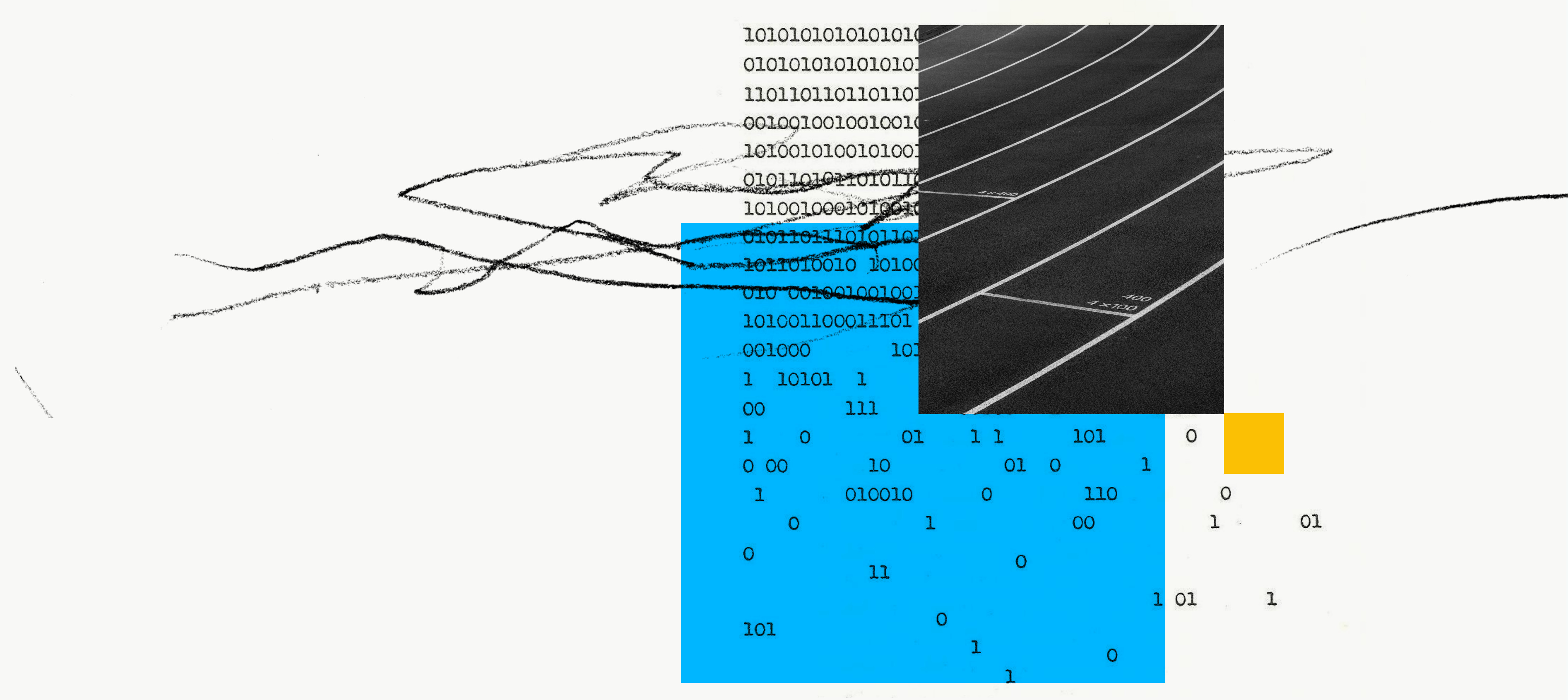


Many respondents report mature capabilities and meaningful reuse, yet most also report rising effort and rising human review. Agentic AI is becoming more trusted and more burdensome at the same time. **The reason is likely that organizations are adding trust controls faster than they’re turning those controls into reusable infrastructure.**

The key operating question is shared: Can the organization make trust requirements predictable enough that they stop arriving as surprises?

CHAPTER 4

TURNING TRUST INTO A COMPETITIVE ADVANTAGE



Respondents often experience trust requirements as a source of speed. Seventy-one percent say those requirements speed deployment of another agent using a similar technical or operating pattern. Across the stages of AI agent deployment, between 68 and 75 percent report a speed benefit. These assessments show perceived operating value; they do not measure how much faster a deployment becomes because requirements are reusable.

Overall, what effect do your organization’s trust requirements have on each of the following stages of AI agent deployment?



Organizations using four or more agent platforms face a sharper version of the Trust Dividend paradox. They report more burden, but they are also more likely to say trust requirements strongly speed several stages of deployment. The same complex environments that feel trust as burden also seem to feel its value as an accelerant.

Overall, what effect do your organization’s trust requirements have on each of the following stages of AI agent deployment?

Getting a first-of-its-kind AI agent from an approved concept into production established before any use-case approval

1-3 agent platforms:

25%

4+ agent platforms:

47%

Deploying another AI agent using a similar technical or operating pattern

1-3 agent platforms:

32%

4+ agent platforms:

46%

Increasing the level of autonomy granted to an existing agent

1-3 agent platforms:

28%

4+ agent platforms:

44%

Trust can be either a bottleneck or an accelerant, depending on the strength and foresight of leadership. It slows organizations when requirements surface late or must be rebuilt from scratch. It speeds organizations when requirements are known, reusable, and verifiable.

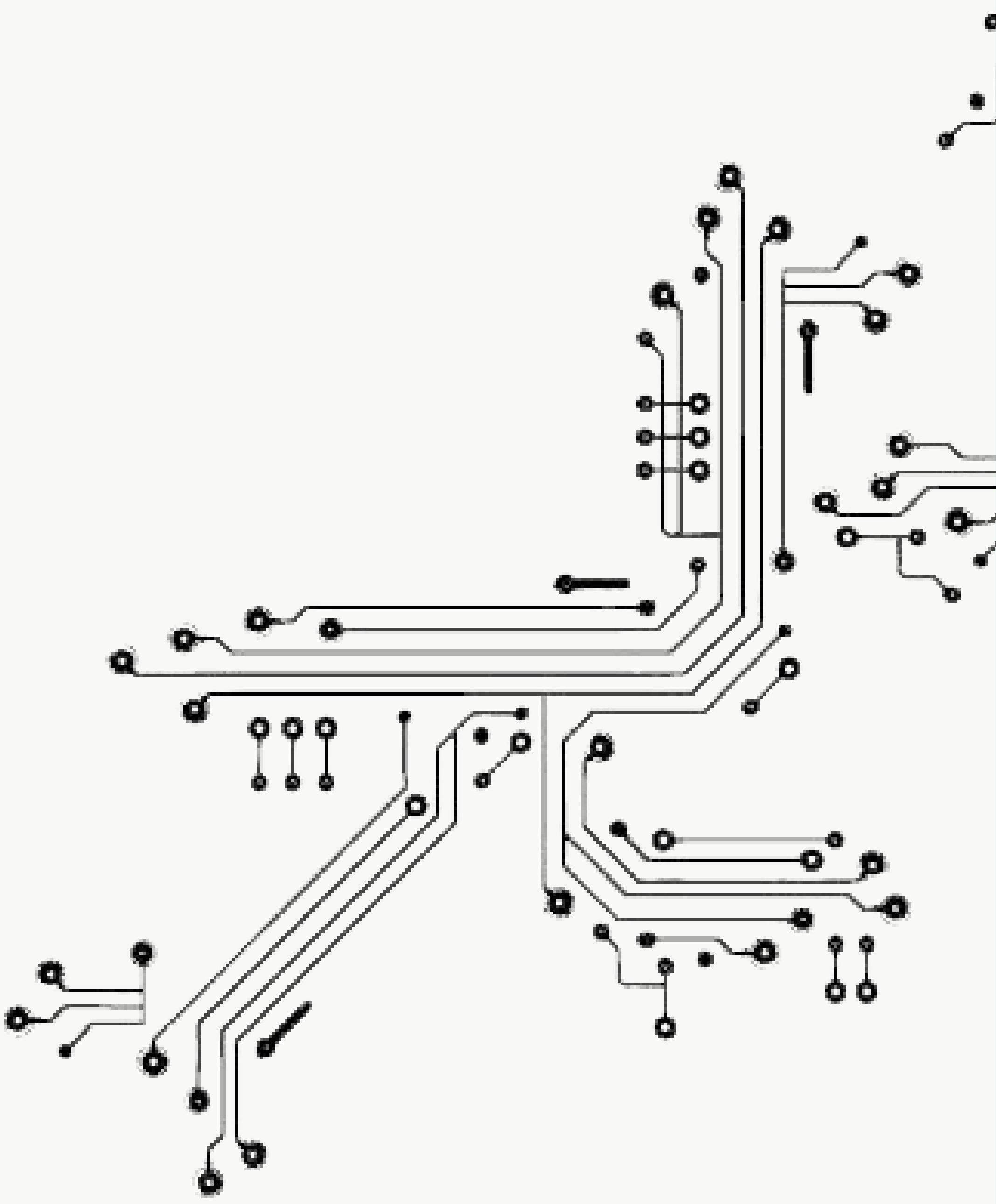
One clear hurdle has emerged for organizations seeking to scale their trust infrastructure: measurement and benchmarking. Safety and reliability incidents are the most commonly tracked agent metric, at 55 percent, followed by output quality and accuracy at 51 percent, and business impact by risk level at 43 percent.

“Because we had solid numbers surrounding the decision, the decision was unanimous and we were able to reach that decision very fast.”

SOFTWARE DEVELOPER

Agent tenure shows what measurement maturity may look like as programs age. Organizations with 12 or more months of production-agent experience are more likely than those with less than 12 months to say they can quantify a distinct contribution from trust capabilities using consistent metrics, 46 percent versus 33 percent. They’re also more likely to track deployment velocity, 44 percent versus 29 percent, and reusable trust infrastructure, 48 percent versus 37 percent.

Where respondents can quantify the Trust Dividend, they describe sizable gains — not marginal ones. Roughly four in five say trust capabilities improved each outcome by at least 10 percent, from deployment speed and approval efficiency to trust reuse, autonomy, and ROI. The effect runs deeper than that threshold suggests: **across all 10 measured outcomes, the average reported improvement exceeds +18%.**



Based on your organization’s measurement or best available estimate, by approximately how much have your trust capabilities improved each of the following AI-agent outcomes?

Average Improvement		Average Improvement	
Deployment speed: Faster movement from trust approval to production	+20%	Autonomy speed: Faster progression from initial deployment to expanded autonomy	+20%
Approval efficiency: Fewer review, approval, exception, or rework cycles	+20%	Safety and reliability: Fewer incidents, rollbacks, or remediation events	+20%
Trust reuse: Greater reuse of controls, testing evidence, or permissions across deployments	+21%	Business impact: Stronger business or customer ROI	+21%
Governance efficiency: Lower governance, testing, and compliance cost	+19%	Output quality: Higher accuracy, acceptance, or quality of agent outputs	+21%
Human oversight efficiency: Less human review, intervention, or override required	+20%	User confidence: Greater end-user or stakeholder confidence in agent performance	+19%

Stronger trust capabilities help organizations match autonomy to the task, move faster, reduce unnecessary approval cycles, and preserve quality as agents do more.

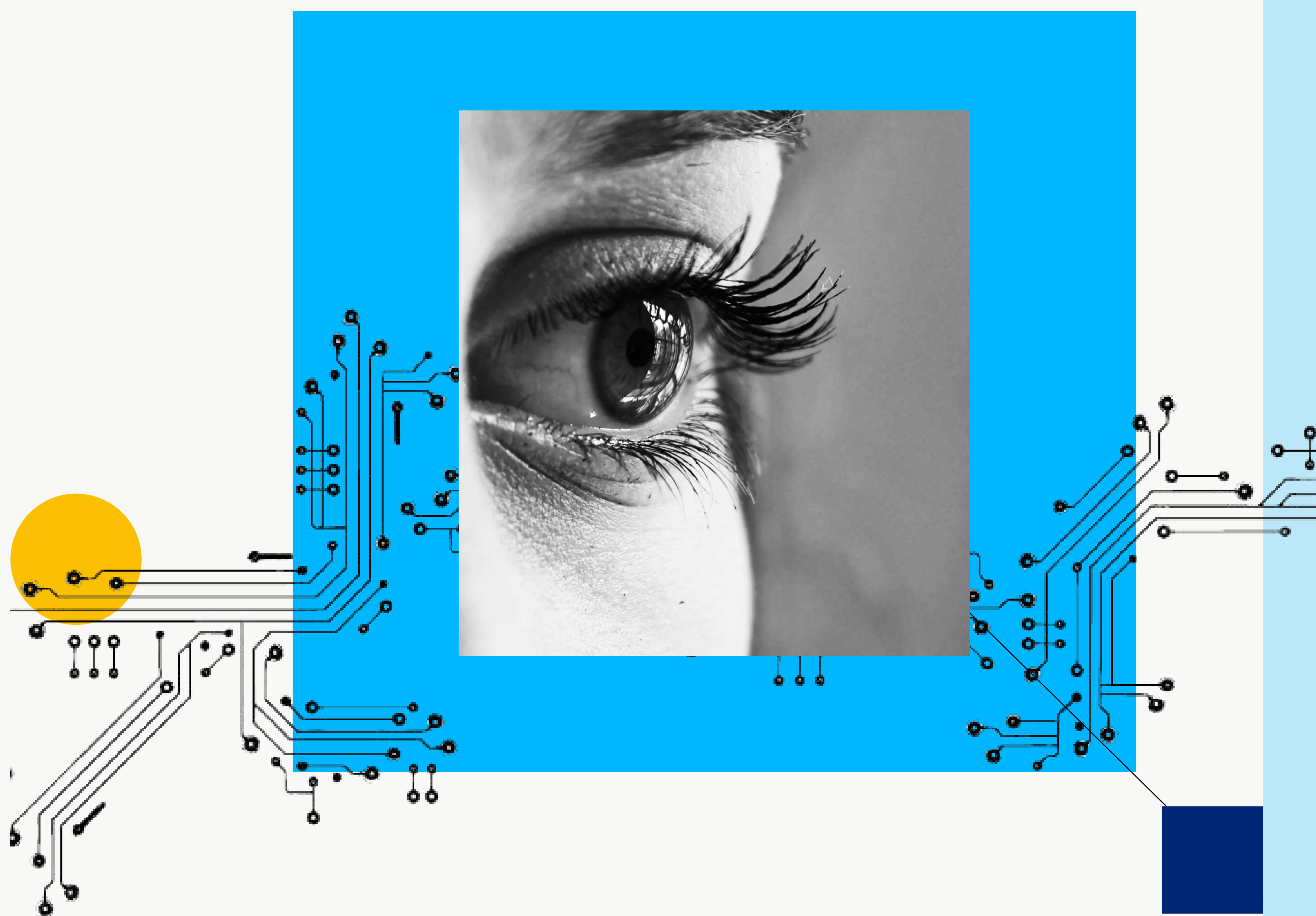
The Trust Dividend compounds over time. A successful agent deployment can produce efficiency once. A reusable trust pattern reduces the work of the next deployment. A tested escalation path can be applied again. A validated permission model can

travel to a similar workflow. A performance threshold can become the basis for the next autonomy decision. Its value is not that it prevents every bad outcome. It’s that the organization knows how to contain the impact, assign accountability, learn from production evidence, and make the next deployment stronger.

Many organizations appear to believe trust improves speed and outcomes, but fewer are measuring the operating variables that would prove it.

CHAPTER 5

FOCUSING HUMAN EXPERTISE WHERE IT MATTERS MOST



The agentic enterprise doesn't remove humans from the system. It changes where human judgment belongs. That distinction matters because the phrase "human in the loop" has become too blunt for agentic AI. It often implies that a person must review or approve each output or action. At agent scale, that model can become a bottleneck disguised as responsibility.

The diary study shows that human oversight takes many different forms. In some workflows, human review protects against irreversible financial action. In others, it preserves professional skepticism, relationship judgment, regulatory accountability, or privacy boundaries. The more mature pattern is not to keep humans everywhere, but to decide where human judgment changes the risk profile most: before money moves, before a customer is contacted, before sensitive data is accessed, before audit work is relied on, or when an exception falls outside the agent's proven competence.

Goldman characterizes this as **mindful friction**: human attention inserted at moments that matter. If the system asks for approval constantly, people stop paying attention. If it asks when an anomaly appears, when a threshold is crossed, or when a sensitive action is proposed, human judgment becomes more valuable. This is the governance equivalent of signal quality. More review time often doesn't translate to meaningfully better quality; better, more precise review does.

Effective managers delegate, set expectations, define boundaries, assess performance, correct course, and decide when to escalate. In the agentic enterprise, every function that uses agents will need some version of that discipline. Once an organization has decided where judgment changes the risk profile for one workflow, that boundary doesn't need to be rebuilt from scratch for the next — it becomes evidence the organization can reuse.

For a sales lead, it may mean deciding which stages of the funnel are appropriate for autonomous outreach. For a CMO, it may mean defining when personalization becomes sensitive. For an operations leader, it may mean identifying which workflow exceptions deserve human intervention. For a CISO, it may mean ensuring that a non-human actor's access can be granted, observed, reduced, and revoked.

“The case-by-case human approval for every low-risk routine communication feels the most cumbersome today. I would not remove human review entirely, but I would consider limiting it to exceptions. Before doing that, I would want several months of results showing at least 98 percent accuracy.”

EDUCATION LEADER

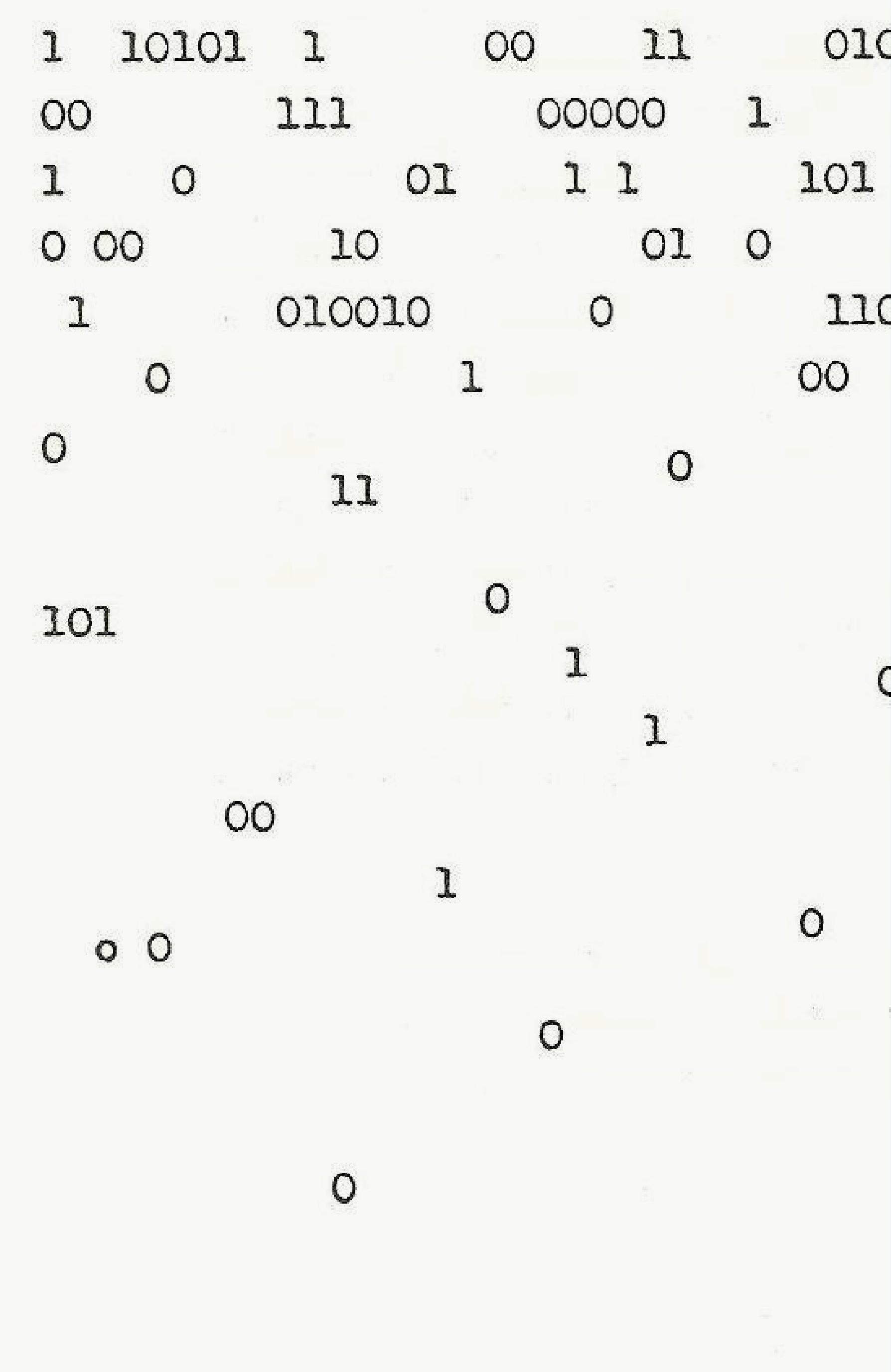
“Once you get a process to agent scale, every subsequent process has to be agent scale as well.”

JOE INZERILLO, PRESIDENT, ENTERPRISE & AI TECHNOLOGY, SALESFORCE

Inzerillo challenges the expectation that an agent must be flawless before it can be trusted. “Just because it’s AI doesn’t mean it’s infallible,” he says. Humans already performing the work make mistakes, too. Leaders need to measure that baseline and establish where an agent improves on it, while accounting for the consequences and scale of its errors.

He draws on the engineering concept of error budgets: explicit limits on the failures a system can tolerate. For agents, that requires judgment about both frequency and severity. A high overall accuracy rate can still conceal an unacceptable failure. Teams need to define which errors can be corrected within the workflow, which require human intervention, and which actions must remain outside the agent’s authority.

Those decisions require shared ownership across business and technology. The team responsible for the workflow must also own its performance thresholds, monitoring, and response when limits are breached. An agent earns latitude through evidence that its errors can be detected, contained, and corrected within those limits. Human judgment sets the boundaries and determines when the evidence supports expanding them. For the next deployment to benefit, those boundaries, the evidence behind them, and the conditions requiring reassessment need to be recorded and available to the team taking responsibility.



CONCLUSION

THE TRUSTWORTHY ENTERPRISE MOVES FASTER

The Trust Dividend isn't a promise that trust makes deploying AI agents easy. The data suggest almost the opposite: as agents gain authority, organizations feel both more capable *and* more burdened. They report reusable trust work, stronger confidence, and perceived gains from better trust features. They also report rising review, rising effort, and frequent pauses when trust requirements surface late. Trust is becoming the inflection point where agentic AI either advances or stalls.

Organizations with longer production-agent tenure are more likely to have platformized trust capabilities and to define trust requirements upfront.

Organizations using four or more platforms report greater governance difficulty, longer authority ramps, and heavier human-review burden, yet they are also more likely to say trust requirements strongly speed key stages of deployment.

Faster deployment without trust creates exposure. Stronger governance without reuse creates process drag. The Trust Dividend appears when organizations convert trust work into reusable infrastructure:

permissions that travel across similar use cases, thresholds that can be revised with production evidence, escalation paths that are known before something goes wrong, and accountability attached to workflows rather than slogans.

The organizations best positioned for the agentic era are those that can explain, measure, and reuse the conditions under which an agent earns more authority. That's the compounding mechanic behind the Trust Dividend. Each deployment leaves behind evidence, controls, and operating judgment that make the next deployment faster and better grounded. ●

Methodology

Atlantic Insights, the marketing research division of *The Atlantic*, and Salesforce conducted a mixed-method study of how enterprises are moving AI agents from experimentation to governed scale.

The research centers on a survey of 300 U.S. business and IT decision-makers at US-based organizations with 300 or more employees. All respondents were director-level or above, and all reported having at least one AI agent operating in live production.

To ground the data in real enterprise practice, the report also draws on a qualitative diary study among leaders using active AI agents, plus in-depth interviews with three Salesforce subject-matter experts: Joe Inzerillo, President, Enterprise and AI Technology; Paula Goldman, EVP, Chief Ethical and Humane Use Officer; and Gagan Gulati, GM of Security.

Atlantic Re:think ×

