

The AI Questions Marketing Leaders Need to Answer Now

A framework for
navigating autonomy,
judgment, legibility,
measurement,
and ownership in
the agentic web.

Last year, Atlantic Insights and Contentful set out to learn what artificial intelligence was doing to the marketer's job. The answer was cheerier than the surrounding panic suggested: The tools absorbed the work marketers liked least, and the teams using it well spent more time on judgment and creativity.

INTRODUCTION

That question has since been overtaken by larger ones. The change that impacts marketing most this year is happening outside the marketing department, in the space between a brand and its customer, where AI systems increasingly research, compare, and narrow the field before a person arrives anywhere a marketer controls.

An answer engine decides what to cite. An assistant assembles a shortlist. An agent compares vendors on a buyer’s behalf. The industry spent a generation optimizing for a person at a search bar. That person now sends a proxy. Increasingly, the first cut is delegated to a machine. Brands may never see the query, enter the consideration set, or know why they were excluded.

Marketing leaders are already treating this as one connected shift. Fifty-one percent say AI-powered search-and-answer engines and AI agents are equally important areas of focus; twenty-nine percent prioritize search-and-answer engines more, and nineteen percent prioritize agents more. Fifty-three percent say they have a defined agentic-commerce strategy and are actively building for it, with another thirty-seven percent exploring agentic commerce without a defined strategy in place. Ninety-one percent agree that the future of marketing increasingly depends on persuading the systems that influence people.

A NOTE ON WHO ANSWERED

Every respondent in this study works at an organization already prioritizing AI; we screened out those that weren’t, because we wanted to learn what happens after the decision to invest in AI has been made. That makes these findings a description of committed organizations rather than of the marketing profession as a whole, which is worth keeping in mind throughout. Among organizations that have already made the AI bet, the problem is no longer whether to invest. It's what happens next.

51%

say AI-powered search and answer engines and AI agents are equally important areas of focus

29%

prioritize search and answer engines more

19%

prioritize agents more

53%

say they have a defined agentic-commerce strategy and are actively building for it

37%

are exploring agentic commerce without a defined strategy in place

This report asks an operational question rather than an attitudinal one. Adopting AI is one thing. Deciding how an organization operates in a market where machines mediate discovery is another, and across 350 organizations we found the same five themes unresolved:

Autonomy. What can our agents do without us?

Judgment. What has to remain human?

Legibility. Can the agentic web understand and verify our brand?

Measurement. What do AI systems say about our brand, and how is that changing?

Ownership. Who is accountable for keeping the system aligned?

Answering these five questions doesn't require a budget. It requires an operating decision on each: what's been decided, who decided it, and when it'll be revisited.

91%

agree that the future of marketing increasingly depends on persuading the systems that influence people.

INTRODUCTION

Some of what follows doesn't have a name yet, which is part of why it goes unmanaged. A team can't assign, schedule, or argue about a cost it has no word for. So each chapter names at least one, from the **the approval tax ↗** to **the accountability map ↗**. coining seven terms to introduce into your company's lexicon. Each is defined where it appears and collected in a glossary at the end. They are the vocabulary for a set of operating problems that have emerged faster than the language for them.

That last part is the point. These are not questions with permanent answers, because the systems they concern are changing faster than the quarterly planning cycle. So each of the five chapters that follow ends with the same four-part framework: **the question** the organization needs to answer, **the risk of deferral** if no one makes the decision, **the operating standard** that constitutes a credible answer, and **the review cadence** that keeps the answer current. Five questions, four parts each: that is **The Clarity Framework**, and every chapter closes by running one question through it.

Used once, it produces decisions. Used on a schedule, it produces an operating model.



01

.....

INTRODUCTION ↗

ONE · AUTONOMY

06

.....

WHAT CAN OUR AGENTS
DO WITHOUT US? ↗

TWO · JUDGMENT

14

.....

WHAT HAS TO REMAIN
HUMAN? ↗

THREE · LEGIBILITY

21

.....

CAN THE AGENTIC WEB
UNDERSTAND AND VERIFY
OUR BRAND? ↗

FOUR · MEASUREMENT

27

.....

WHAT DO AI SYSTEMS SAY
ABOUT OUR BRAND, AND HOW
IS THAT CHANGING? ↗

FIVE · OWNERSHIP

33

.....

WHO IS ACCOUNTABLE
FOR KEEPING THE SYSTEM
ALIGNED? ↗

39

.....

CONCLUSION ↗

41

.....

METHODOLOGY ↗

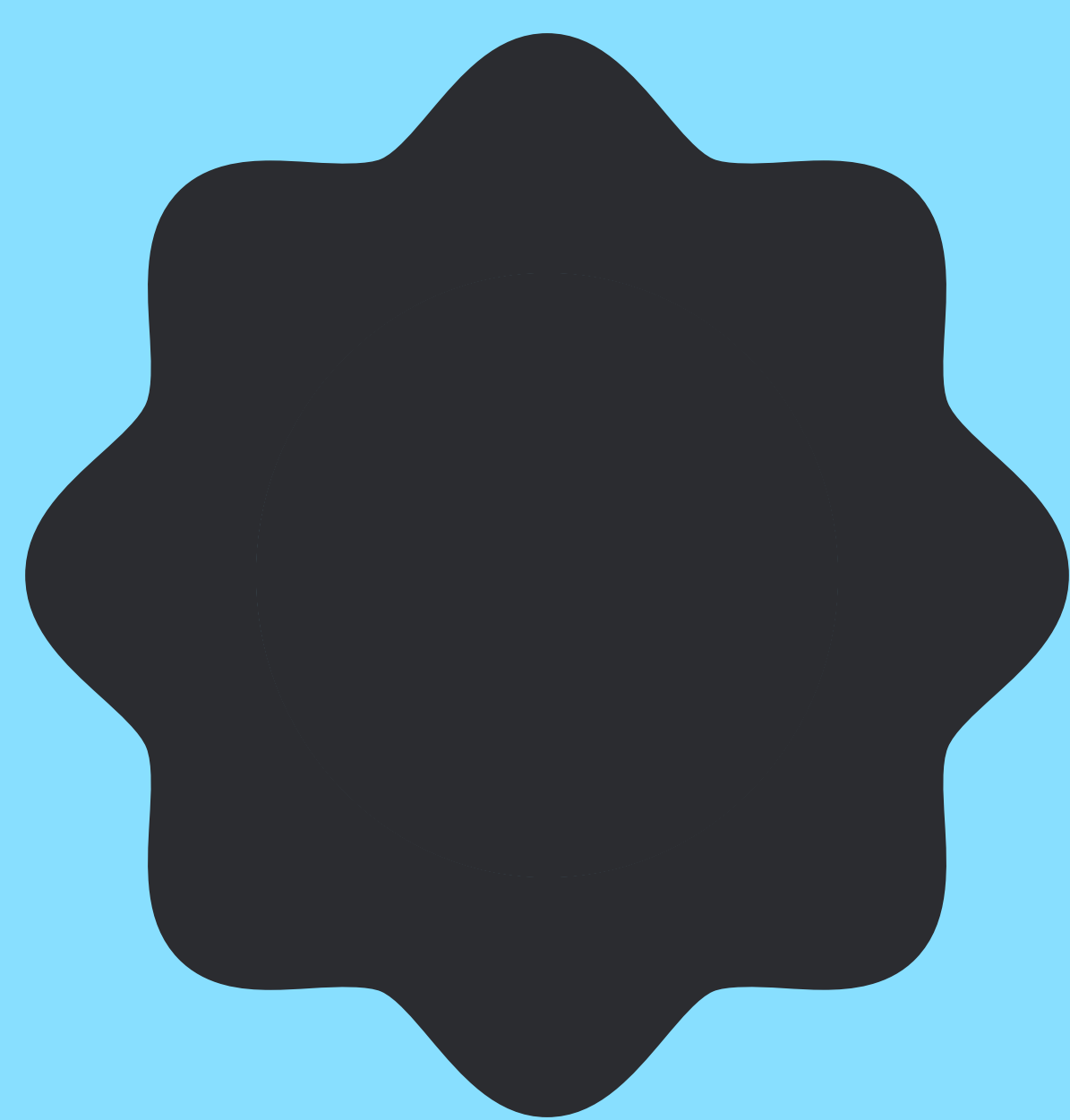
42

.....

GLOSSARY ↗

What Can Our Agents Do Without Us?

An agent that answers a customer question, updates a product page, or adjusts a live campaign is doing something the marketing department has never had to govern before: acting at the boundary between the brand and its market, with nobody in the room.



KEY TERM

THE APPROVAL TAX

what routine work costs while it waits for a person.



KEY TERM

DECISION BOUNDARY

the written line between agent, reviewer, and decider.

The autonomy question arrives looking like an internal governance matter. But it resolves into an external one, because whatever an agent does unsupervised is what a customer, or another machine, experiences as the brand.

3.6

Agents are working in an average of 3.6 marketing functions per organization.

Most organizations are already there. Agents are working in an average of 3.6 marketing functions per organization. Search, SEO, and AEO management lead at forty-eight percent, followed by customer service and post-purchase engagement at forty-six, and campaign personalization and segmentation at forty-four. Only two percent report no agent use anywhere in the organization.

What hasn't been settled is how far they can go alone. Seventy percent of teams already describe agents as trusted with important work or being scaled across the organization.

Forty-three percent would let an agent act autonomously across most marketing tasks, subject to periodic review. Thirty-four percent confine agents to low-stakes decisions. Twenty-one percent require an agent to propose every action and wait for approval before anything executes.



Agent Functions:

46%

customer service and post-purchase engagement

44%

campaign personalization and segmentation

48%

search, SEO and AEO management

2%

report no agent use anywhere



“How do I enforce workflow to ensure that anything that has been touched or used by AI has been reviewed by a human before it goes out the door?

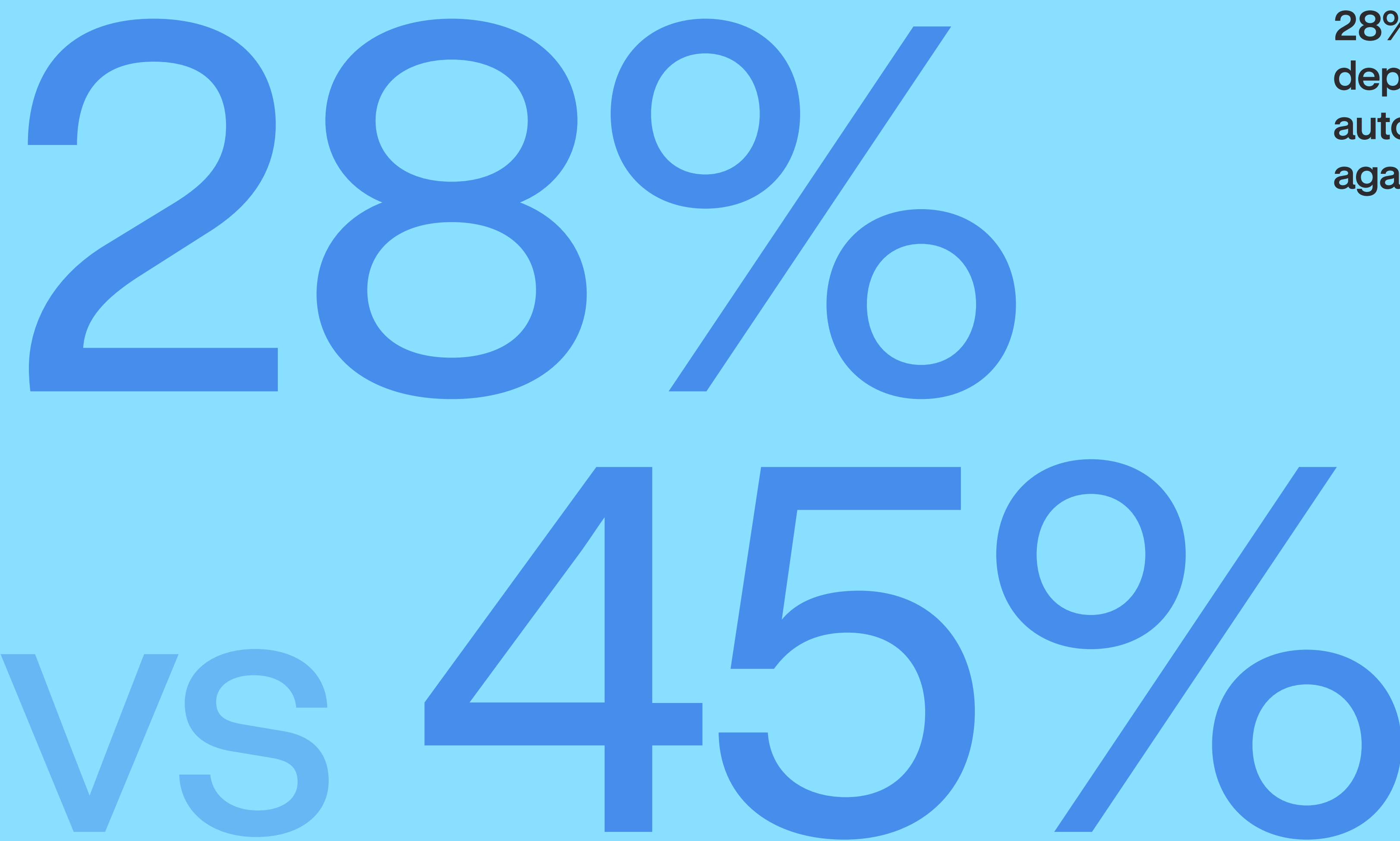
— SARA SULLIVAN, SVP OF SOLUTION ENGINEERING, CONTENTFUL

Some checkpoints are deliberate, often set by people outside marketing. But caution has a cost. An **approval tax** ↗ is the time and value lost when an agent waits for a human review no one has decided is necessary. Some reviews earn their keep — a false claim caught before it reaches a customer is, of course, time well spent. The tax comes from approvals that exist because no one drew the line. And it remains largely invisible: Organizations measure the work that gets done, not the work waiting for permission.

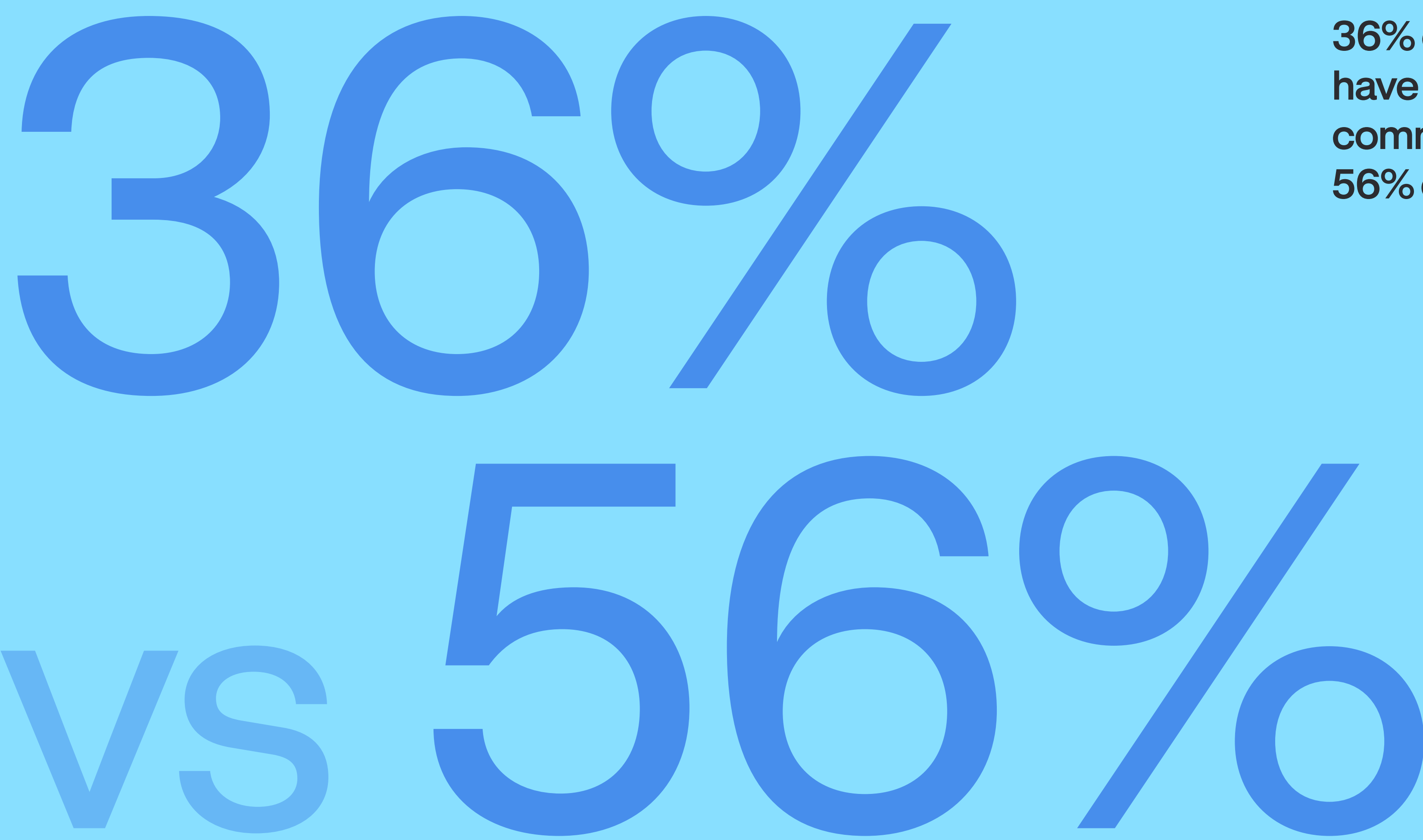
Small teams, large teams.

Autonomy is not evenly distributed, and department size predicts it more reliably than seniority. Twenty-eight percent of small marketing departments would extend autonomy across most tasks, against forty-five percent of larger ones. Thirty-six percent of small departments have a defined agentic-commerce strategy, against fifty-six percent of large ones. And thirty-eight percent of small departments feel extremely well-equipped for an AI-mediated web, against fifty-eight percent of large ones.

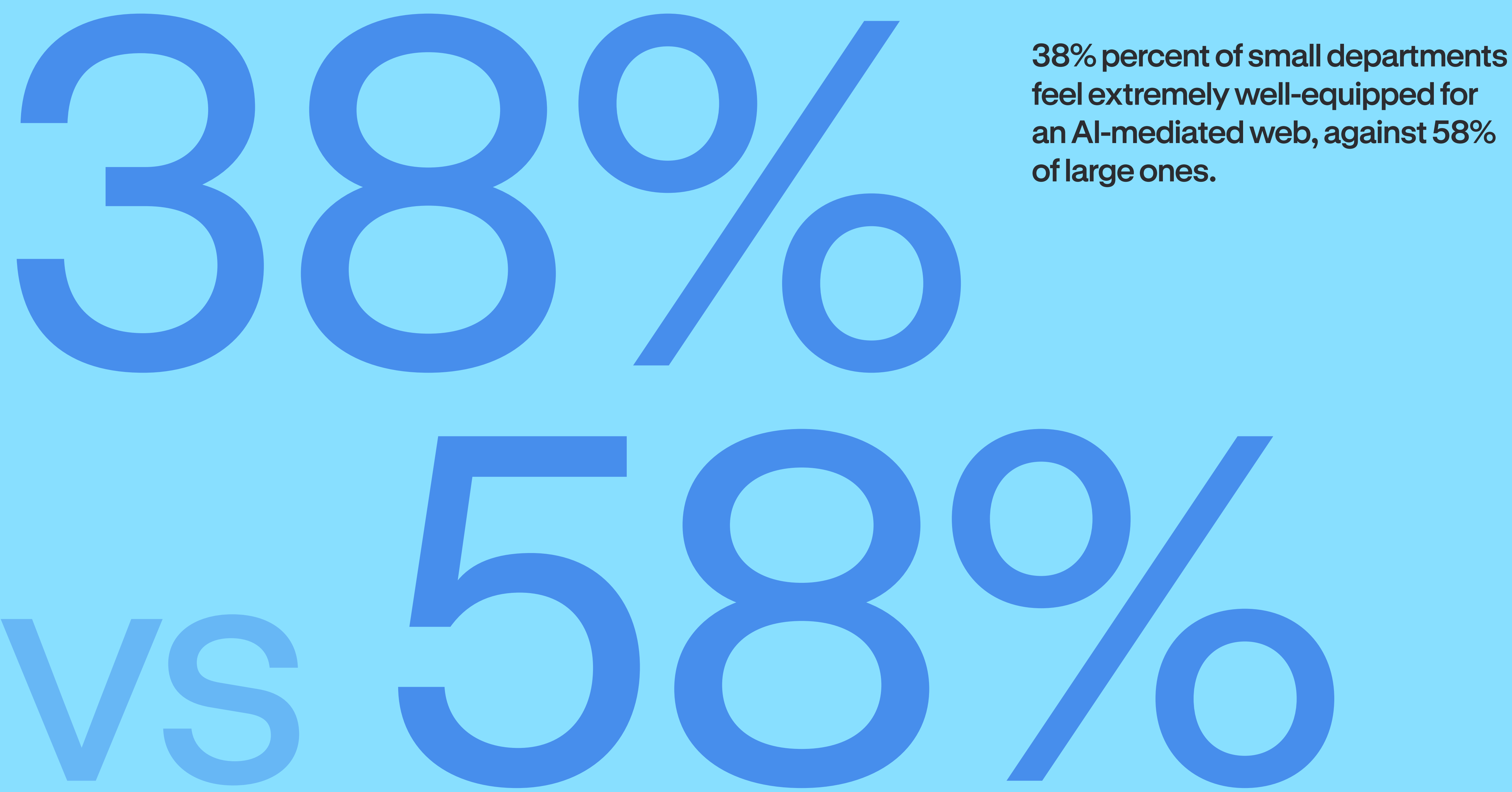
All three gaps hold up statistically.



28% percent of small marketing departments would extend autonomy across most tasks, against 45% of larger ones.



36% of small departments have a defined agentic-commerce strategy, against 56% of large ones.



What sits behind them is resourcing. Leaders of large marketing departments are three and a half times more likely to be working with a \$500M company's budget behind them; leaders of small departments are more likely to sit at companies with under \$200M in revenue, and are fifty-six percent less likely to control a large marketing budget themselves. The readiness gap tracks the money.

That reframes what the numbers are measuring. A small department extending autonomy across fewer tasks is not a more cautious department it's one with less room to absorb a mistake, fewer people to write the boundary, and less budget to buy the systems that would monitor it.

Small departments here means fifty or fewer people, which describes fifteen percent of this sample (n=53). If you run one benchmark, the one that matters is other small teams. The resourcing gap is real, and measuring against enterprise marketing organizations produces a readiness score you can't act on.

Charlie Bell, who leads Contentful's Solutions Engineering in EMEA, Australia, and New Zealand, expects some version of that cost to be permanent. However capable agents become, he expects them to be trusted about ninety percent of the way and no further, with the ninety percent itself a moving target: ninety percent of what the job requires now, then ninety percent of what it requires next year. The last ten percent belongs to whoever puts their name on the work.

Read that as an argument for precision rather than restraint. A cost an organization pays every day is worth deciding on deliberately, and the way to reduce it is to know exactly which decisions actually need a human to sign off.

Most organizations have already drawn that line instinctively, along the boundary between reading and writing. In our qualitative research, executives were markedly more comfortable letting agents find patterns and surface opportunities than letting them create and publish. The hesitation clustered around permissions, auditability, and control rather than capability.

The instinct is a good one. It fails when it goes unwritten. These inconsistencies reach the market as brand behavior long before anyone traces them back to a decision nobody made. **What the room needs is a [decision boundary](#) ↗: a written line between work an agent can do, work a person reviews, and work a person must decide.**

90%

Charlie Bell, who leads Contentful's Solutions Engineering in EMEA, Australia, and New Zealand, expects agents to be trusted about 90% of the way and no further.

The Clarity Framework:

Autonomy

THE QUESTION

What can our agents do without us? Bring three real examples of work an agent is doing now, or is being considered for. Not hypotheticals.

-
- What can the agent decide,** and what can it change?
 - What information and systems** can it reach?
 - Which actions need approval first,** and which can be reviewed after?
 - Who is accountable** if it sends the wrong message or acts on bad data?
 - Would another team** sort the same action the same way?

Run this at the level where the work actually happens. A hundred-person marketing organization should do it in functional groups, not as a single leadership exercise, and the executive who sets risk tolerance should see all the answers side by side.

The Risk Of Deferral

Without a written boundary, autonomy is set by whoever is improvising that day. Two teams sort the same action differently, the approval tax accumulates in the places nobody measures, and the first time the organization discovers its actual policy is during an incident.

The Operating Standard

A permission structure that reflects the cost of being wrong, where the executive sets the risk tolerance. Every agent has a named owner, defined data permissions, and explicit limits on what it can change. Low-risk, reversible work proceeds with periodic review; anything touching brand voice, claims, spend, or targeting gets human review first. Every action leaves a record of who approved it, and every workflow has an off switch.

Three tiers, written down.

The Review Cadence

Re-sort the three tiers quarterly, and whenever a new agent, model, or integration enters the stack. Track how often work sits in the approval queue and why. That number is the approval tax, and it is the evidence for loosening or tightening the boundary.

What Has to Remain Human?

As machines take on more of what passes between a brand and its customer, the human contribution focuses on direction. Deciding what the brand should say. Judging whether a system has represented it correctly. Recognizing that an answer is wrong before a customer acts on it.



KEY TERM

JUDGMENT DEBT

the future cost of automating the work that taught people to decide well.

The scarce thing in an agentic organization is not execution capacity. It is the judgment to know whether what was executed is any good.

That work gets more valuable as machine-mediated interactions multiply. A marketer used to approve a piece of work and see it run once. Now they approve a rule, and agents apply it across every exchange that follows.

Marketers already know this, and the data shows them protecting it. Creative ideas and strategic decisions are tied at the top of that list, followed by brand voice, then customer-facing copy and performance insights. All five sit within four points of one another. No single task is markedly more protected than the others, because what marketers are protecting is judgment as a category.

Human creativity and judgment aren't a matter of spotting style. They are the ability to hold an objective and evaluate work against it, which is exactly the capability that operational work used to build.

Which makes the restructuring data worth reading twice. Ninety-six percent of marketing teams report at least one structural change in the past year because of AI. Most of it is investment in capacity. Forty-six percent folded automated systems or “digital workers” into team workflows, and forty-three percent added roles focused on AI or automation. Fifty-three percent increased AI training or enablement, though instruction in a tool is not the same as learning to evaluate an idea with human judgment, which is the distinction this chapter draws.

A smaller set of decisions runs the other way, and they are a different kind of decision. Twenty-one percent eliminated certain roles, and twenty-five percent reduced or paused entry-level hiring.

25%

of marketing teams reduced or paused entry-level hiring in the past year because of AI.

That last number is the one to sit with, because its consequences arrive late. Judgment about quality, brand fit, and strategic direction has historically been built by doing the operational work (often done by junior employees) that agents now absorb. The tagging. The trafficking. The fourteenth revision of an A/B-tested subject line, which nobody enjoyed but everybody learned from. Set the automation data beside the trust data and the two lists are near-perfect complements: Agents have been given the work that reaches the customer at volume, and kept away from the work that decides what reaches them. That allocation is one that removes the training grounds for judgment.

Two lines are moving in opposite directions. Human judgment becomes more important as machines assume more of the interaction, not less, because fewer people are now carrying more consequential decisions. Meanwhile, the same automation that raises the value of that judgment is removing the experiences that employees need to hone their discernment. **Judgment debt ↗ is the distance between those lines: the future cost of automating the work that taught people to decide well, during the exact period the organization comes to depend on that skill most.** The debt is incurred at the entry level and paid at the senior one by a different set of managers than the ones who took it on.

“To write effectively, to communicate effectively, you need to know what your objective is, and you need to be able to measure its impact... not just say, are there em dashes, but rather, does this prove my objective? Did more people buy my product? Did more people share this article?”

— GABRIEL DILLON, GO-TO-MARKET LEAD, PERSONALIZATION, CONTENTFUL

Where Marketers Trust AI Least

(Q26, n=350)



It doesn't appear in this year's efficiency numbers. Judgment debt appears when the organization needs someone who can look at a campaign and say why it is wrong, only to realize the roles that used to produce that judgment were automated in 2026.

None of this is invisible to the people running these teams — the roles just aren't being funded. Which brings this section back to where it started. Seventy-nine percent of marketing leaders agree the teams winning in the AI era will have the best operations, not the strongest creative. Both things are being said at once. The reconciliation is that operations gets budget, headcount, and tooling, while judgment development is treated as something the organization already has. Only one of those is being built deliberately.

Most routine tasks aren't worth keeping for their own sake. The risk is removing them without deliberately replacing the feedback, customer exposure, and senior coaching that came bundled with them. When integrating AI, organizations should be mindful of not removing the apprenticeship layer businesses depend on to create their next generation of experts.

Sullivan draws the line at origination. "Do you really want AI to generate your next great marketing idea?" she says. "They're probabilistic LLMs, so everything they're going to produce is something that has existed before, and we believe that the best ideas come from humans."

79%

of marketing leaders agree the teams winning in the AI era will have the best operations, not the strongest creative.

The Clarity Framework:

Judgment

THE QUESTION

What has to remain human? Name the work that used to teach junior employees how the brand, the customer, and the business actually behave.

Which of those tasks are now automated?

What has replaced the learning they provided?

How does a junior person learn to tell strong work from merely acceptable work?

Who explains why a piece of work is right for the brand?

Which customer-facing decisions stay human?

How do we measure whether people's decisions are improving?

The Risk Of Deferral

Judgment debt compounds silently. An organization that removes the bottom rung while betting its future on the top one is running a deficit it won't feel for years and won't be able to close in one. We can look at this through the lens of 2025 research by Microsoft Research and Carnegie Mellon researchers that found that high confidence in GenAI is associated with less critical-thinking effort, while self-confidence predicts more. When organizations automate the work through which people learn to verify, compare, diagnose, and defend a decision, they risk replacing deliberate practice with acceptance of a system's first answer. AI does not make that outcome inevitable, but it shifts the responsibility of that learning to management rather than naturally being a by-product of junior work.

The Operating Standard

Judgment development sits inside the AI strategy. Leaders can identify which tasks were automated and what learning accompanied them. Junior marketers still get structured exposure to edge cases, customer context, and real tradeoffs. Senior marketers explain why work is strong, weak, or off-brand instead of only approving it. A stated policy says where AI does and doesn't touch customer-facing work. Training is judged by the decisions people can make, not courses completed.

The Review Cadence

Twice a year, alongside workforce planning: which tasks were automated this year, what learning went with them, and what replaced it. Track the decisions your people can make now that they couldn't a year ago. That is the only measure of whether the pipeline is working.

Can The Agentic Web Understand And Verify Our Brand?

Eighty-five percent of marketing leaders believe AI summarization will make most brands sound the same. The instinct is to answer that by publishing more, which adds to the problem it is meant to solve.



KEY TERM

CANONICAL BRAND SOURCE

the one governed place those facts are maintained.



KEY TERM

SIGNAL INTEGRITY

whether a brand’s facts hold up across every source a system can reach.

Sameness costs something, and the cost isn't aesthetic. "Human beings are exceptionally good at spotting patterns," Bell says, and audiences have learned this one: the structure, format, and phrasing of copy a machine produced. Once they recognize it, "the value curve for that content has dropped." The volume strategy defeats itself, and it defeats itself twice over, because the same undifferentiated material is what an answer engine reads when it decides what your brand is.

The systems doing the summarizing assemble their understanding of a brand from websites, product data, metadata, reviews, third-party coverage, and citations. When those sources disagree, the system resolves the conflict on its own terms. Sullivan describes the check she would want running before anything publishes:

"Do we have other messages already out in our digital properties that say something different? And now all of a sudden we're going to confuse these answer engines because we have two conflicting messages."

Signal integrity ↗ is the shorthand for that consistency: the degree to which a brand's facts, claims, and proof hold up across every source an AI system can reach. Asked what actually distinguishes a brand once AI is doing the summarizing, thirty-seven percent of executives point to the structure and clarity of content. Twenty percent point to third-party validation and authority. The content itself ranks fourth, at seventeen.

Both of those numbers conceal a split. Australian respondents put far more weight on structure and clarity than their counterparts elsewhere: Fifty-four percent named it, against thirty-five percent in the United States, and thirty-two percent in the United Kingdom. On third-party validation, the divide runs by size rather than geography. Twenty-two percent of large marketing departments name it, against nine percent of small ones.

95%

of respondents agree that brands with strong, well-codified identities will widen their lead in an AI-mediated world.

“Content needs to be structured and well organized. That means that there’s a declarative understanding of what the content is and what it does, so that agents can then access that information and know not just what are the words within the content, but how that content can be applied.”

— GABRIEL DILLON, GO-TO-MARKET LEAD, PERSONALIZATION, CONTENTFUL

Read together, they say the route to becoming distinguishable is not singular. One approach starts at the source: Make the brand's own account of itself clear enough, and structured enough, that a system has something unambiguous to read. The other starts outside the brand: Make sure enough credible third parties describe it the same way, so a system finds corroboration wherever it looks.

A sound strategy incorporates both.

Read the top of that ordering again. Volume no longer creates an advantage in AI visibility, because AI has made content cheap to produce. What creates advantage is a clear, well-structured, well-corroborated account of what the company actually is, maintained somewhere a system can reach it. **That account is a canonical brand source ↗: the one governed place where the organization's facts, claims, product detail, and approved language live, and the place every other system is required to draw from.** Signal integrity is the condition; the canonical source is what produces it.

The building has started, but no single action has become standard practice. No optimization action in this study is prioritized by more than thirty-four percent of respondents. Keeping brand information consistent across third-party sources and maintaining dedicated FAQ or Q&A pages both reach thirty-four percent. Schema markup, the structured tags that tell a machine what a page’s content means and not just what it says, reaches thirty-one. An MCP server or similar agent-facing infrastructure, which lets an AI system query a company’s information directly rather than inferring it from published pages, reaches twenty-nine.

Stated priorities have already moved. AEO and GEO, answer engine optimization and generative engine optimization, are now prioritized over SEO, fifty-five percent to forty-five. This is a shift in emphasis, not a succession. AEO and SEO work alongside each other: Answer engines lean on classic search infrastructure for information, while people lean increasingly on the answer engines. Content has to win twice. What’s published hasn’t caught up with what’s prioritized.

The qualitative research makes a sequencing point the survey can't.

Asked where AI exerts its greatest influence on the buying process, the 22 executives in our conversational session put it at discovery (seventy-five percent), information gathering (sixty-five), and recommendations (sixty). Brand evaluation drew forty. Base size is small; treat as directional context rather than a benchmark.

(Q26, n=22)



The Clarity Framework:

Legibility

THE QUESTION

Can the agentic web under-stand and verify our brand? Take one product, proposition, or customer promise, and trace it through the places an AI system might look.

Where is the official version maintained, and is it current?

Is it organized so that a system can understand and reuse it?

Does the same claim hold on the website, in product data, in press coverage, on review sites like G2, and in community sources like Reddit?

What happens when two sources disagree?

What evidence supports it?

Could a system explain how we differ from our closest competitor, or would it fall back on category language?

The Risk Of Deferral

When sources conflict, the system picks.
A brand that hasn’t agreed internally on what is accurate has effectively delegated that decision to whichever source is easiest to parse, often a competitor’s description of the category or an outdated page nobody owns.

The Operating Standard

A single governed source of brand facts, covering company information, product detail, claims, evidence, and approved language, that every downstream system draws from. Content is structured, tagged, current, and reusable, with product and marketing agreed on which facts stay consistent. That version holds up across third-party sources too. New content is checked for accuracy, brand, and metadata before it is published.

The Review Cadence

Test on a schedule whether a machine can find you, understand you, and tell you apart. When it can’t, name the failing source rather than rewriting the answer. Re-audit third-party and review sources at least twice a year, and after any repositioning, pricing change, or product launch.

What Do AI Systems Say About Our Brand, And How Is That Changing?

Eighty-three percent of marketing leaders in our quantitative survey call showing up accurately in AI systems a top or high priority for the next twelve months, but just more than half of marketing teams (fifty-five percent) are regularly measuring how they appear in the agentic web.



KEY TERM

**REPRESENTATION
DRIFT**

the slow change in how AI systems describe a brand.

Far fewer are set up to see **representation drift — a change over time in how AI systems describe, compare, or recommend a brand.**

The difference is between checking and comparing. A check tells you what a system said today, but you'll need a definitive baseline to track drift. This can be accomplished by asking the same buyer questions, in the same way, with last quarter's answers on file for comparison.

So the tooling question is largely settled. Answer-engine measurement exists, and a majority of these teams are using some form of it. What many still lack is the standardized attribution and accountability model they spent decades building around search: an agreed set of buyer questions, a fixed cadence, a baseline to compare against, and a route from a change in the answer back to the information that caused it.

That distinction matters more here than in most measurement arguments, because the failure is quiet. A product fact goes out of date. An important distinction disappears from a summary.

A competitor becomes the source the system cites most often. None of this registers as a drop on a dashboard you already watch, because none of it happens on a channel you already measure. The effect arrives later, as pipeline you can't attribute, from buyers you never identified, who acted on a description of your brand that nobody at your company approved.

That is what representation drift looks like from the inside — nothing visibly wrong, until the pipeline arrives short.

Bell can describe the far end of that failure because he was on the buying side of it. Replacing the family car after his eldest left for university, he put two models into an assistant and asked for the differences. It walked him through the specifications, then offered to work out cost of ownership, then took him through fuel, insurance, maintenance, and depreciation as he fed it his own numbers, and arrived at a five-year saving of several thousand dollars for one car over the other. Whether the figure was correct, he says, is beside the point.

“The question is not whether the data that it gave me was accurate... The question is whether it is convincing enough for it to make a decision. And in this case, one manufacturer got a sale, one manufacturer didn't get a sale.”

— CHARLIE BELL, SENIOR DIRECTOR, SOLUTIONS ENGINEERING, CONTENTFUL

92%

of US respondents have measured how AI systems describe their brand at least once.

While in Australia the figure is

72%

Neither manufacturer will ever learn the comparison happened. The one that lost has no record of the buyer, the query, or the reasoning, and nothing in its dashboards will move in a way that points to any of it.

The practices that would catch this are still uneven. Thirty-three percent of respondents track recurring buyer questions to fill content gaps. Thirty-one percent have built or acquired a tool to measure agent visibility. Twenty-four percent audit brand mentions in agent outputs. These are reasonable things to do, and none of them are being done by even half the sample. That is a statement about operating discipline rather than about what's available to buy.

There's some notable geographical texture here too. Ninety-two percent of US respondents have measured how AI systems describe their brand at least once. In Australia the figure is seventy-two percent, a significant drop in measurement. Australian marketers are also the most worried about the thing measurement is for. Forty percent name being inaccurately represented by AI systems as their single greatest concern, against twenty-two percent in a place like the UK, another significant gap.

That gap is the distance between how acutely a market feels the problem and how routine the checking has become. The teams most concerned about being described wrongly are the least likely to hold a record of how AI systems describe them. That is why a regular measurement cadence matters: Without a record, teams cannot see what has changed or know what to fix.

Sullivan reduces the work to three questions. First, do you understand what the systems are saying, and whether it's accurate? Second, can you act on what you find without rebuilding the same thing in a dozen places? Third, before anything new goes out, has somebody checked that it's on brand, correctly tagged, and consistent with what's already published? The order matters. Most organizations in this study have started somewhere in the middle.

Among the 22 executives in our conversational session ...

(Q26, n=22)

... ninety-one percent believe visibility in AI recommendations will become more important, fifty percent prioritize it today, and forty-five percent monitor it even informally. Base size is small; treat as directional context rather than a benchmark. The gap between how leaders describe themselves in a survey and how they describe their problems in conversation is itself the finding.



The Clarity Framework:

Measurement

THE QUESTION

What do AI systems say about our brand, and how is that changing? Ask the owners of AI visibility to show what the major systems said about the brand last month. A live prompt is not a substitute for a record.

Which ten questions would a real buyer ask before choosing us?

How did the systems answer them?

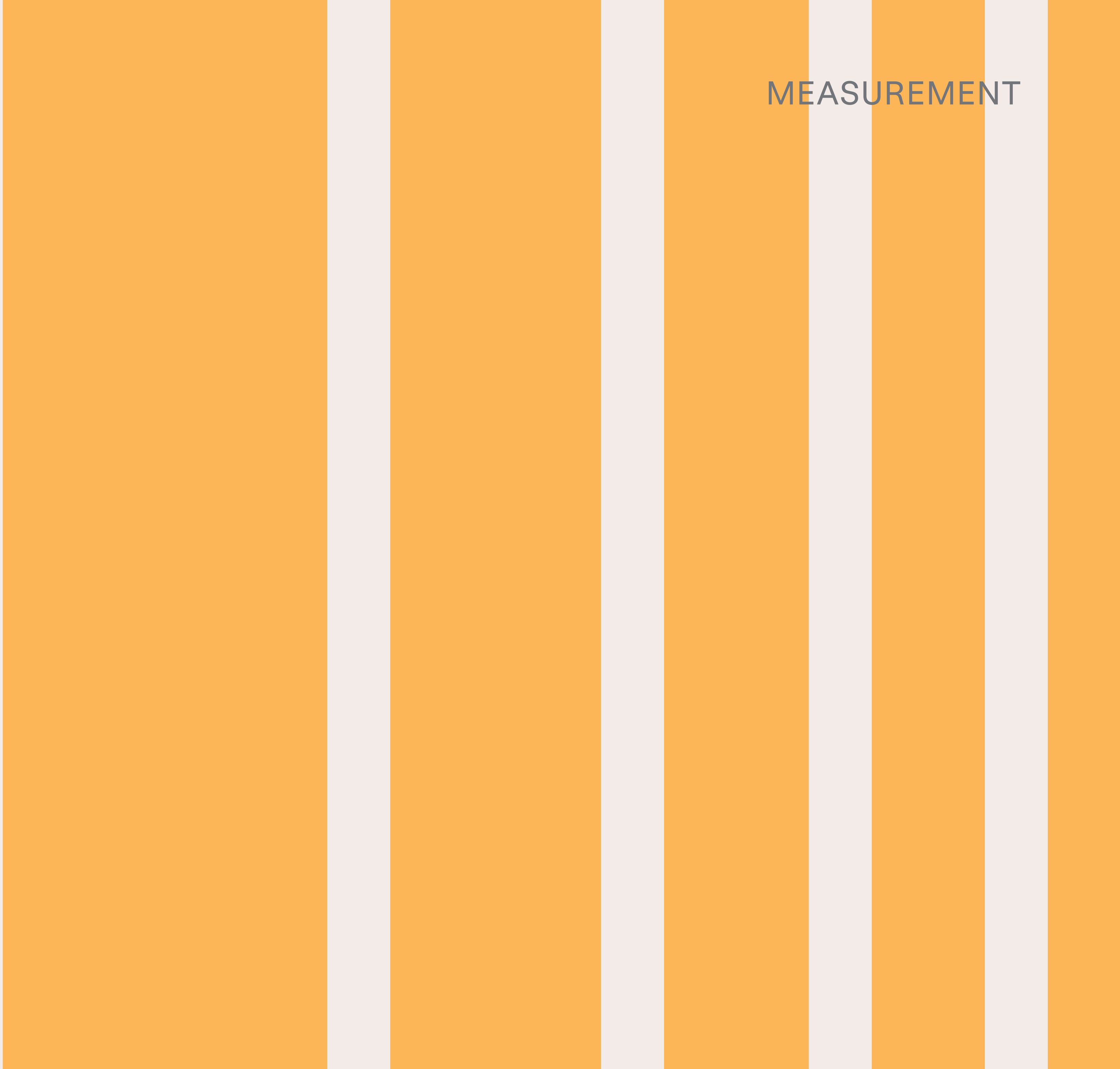
Which competitors appeared, and which sources got cited?

What was inaccurate, missing, outdated, or generic?

How does that compare with the last time we ran it?

Who owns fixing the source, rather than the answer?





The Risk Of Deferral

Without a baseline, every check is a snapshot, and drift becomes visible only after it has cost something. Without a route from answer back to source, the team fixes symptoms, rewriting a page or correcting a single output, while the input that caused the error stays in place and reproduces it.

The Operating Standard

The CMO owns the outcome, with a named lead for the measurement itself. A fixed set of buyer questions runs on a schedule across the systems your customers actually use, recording whether the brand appears, how it’s described, who else shows up and what gets cited. Wrong answers get traced to a cause, the fix gets checked, and the correction workflow reaches content, product data, structured information, and third-party sources. The executive reviews the trend, not the snapshot.

The Review Cadence

Monthly at minimum for the core question set, reviewed alongside the rest of the business, not in a separate AI meeting. Refresh the question set twice a year so it still reflects how buyers actually decide, and re-baseline after any major product, pricing, or positioning change.

Who Is Accountable For Keeping The System Aligned?

Ask marketing leaders what worries them most about AI and discovery, and internal skills and ownership finishes last, at four percent.



KEY TERM

**ACCOUNTABILITY
MAP**

who owns brand guidance, product truth, data access, and the fix. Consider this akin to a RACI matrix (Responsible, Accountable, Consulted, Informed).

Ownership is the mechanism that determines whether the other four questions can be answered at all, and whether the answers hold up a year later.

Rightly, their attention goes first to the failures they can see in the market: being inaccurately represented by AI systems (thirty-one percent), losing direct traffic and owned engagement (twenty-two percent), and losing visibility in AI-driven discovery (twenty-one percent). Falling behind competitors in optimization follows at thirteen percent, while nine percent don't see AI's influence on brand discovery as a concern yet.

Ownership finishes last because it isn't the form the problem takes. Nobody's quarter goes badly because of an org chart. It goes badly because a claim was wrong, a source was stale, or a correction never reached the system that needed it. Every one of those failures traces back to a responsibility nobody had assigned, and every one of them reaches a customer through a machine before anyone inside the company sees it.

Ownership is the mechanism that determines whether the other four questions can be answered at all, and whether the answers hold up a year later.

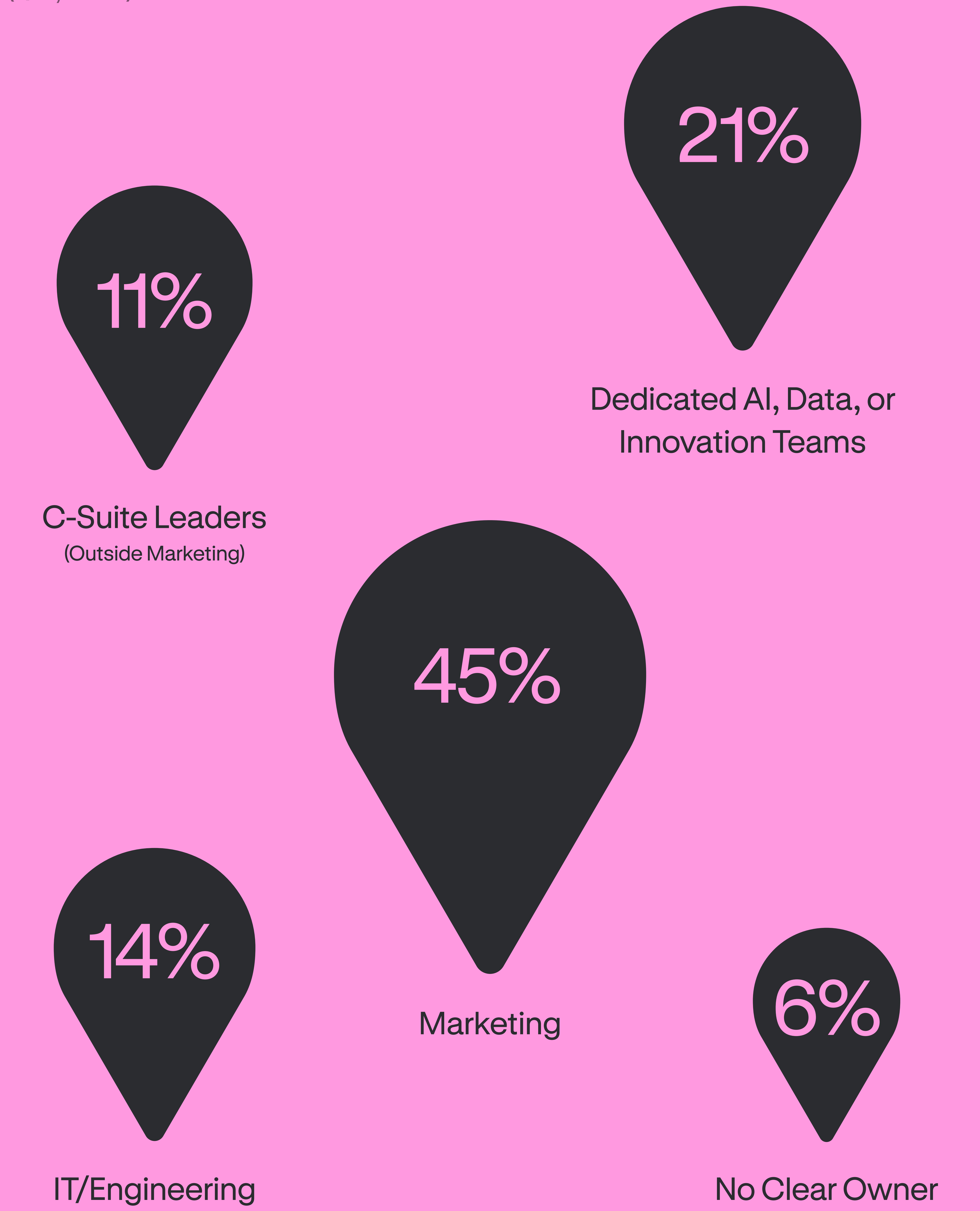
The work currently sits in different places.

Marketing holds primary responsibility for brand visibility in AI systems at forty-five percent of organizations. A dedicated AI, data, or innovation team holds it at twenty-one percent, IT or engineering at fourteen, and C-suite leaders outside marketing at eleven. Six percent report no clear owner, or shared ownership without a clear owner. That is a description of where the work has landed, not an argument for where it belongs.

There is a reason for the dispersion. Visibility in AI systems touches content, product information, pricing, thought leadership, third-party coverage and, awkwardly, the language competitors use about themselves. It doesn't sit under any existing specialist, which is why it keeps landing on a committee.

Who Owns The Responsibility?

(Q26, n=350)



Emerging AI Roles



Org charts are absorbing that ambiguity rather than resolving it. Asked which emerging AI roles exist on their teams or are under consideration, half name an AI or agent strategist; governance, ethics, or disclosure leads reach forty-two percent, and AEO managers, agent-operations orchestrators, and AI brand-safety officers each reach thirty-seven. New titles are proliferating faster than decision rights are being assigned.

No single function controls all the inputs, which is why the answer is to separate accountability from responsibility.

One executive is accountable for the brand outcome and answers to the C-suite for it. A small group of named functional owners is responsible for the parts they control: brand and content, product, and data and technology. Each owner has explicit decision rights, a shared source of truth, and a defined escalation path. The executive doesn't need to know everything at all times; the group exists so that answers are available on demand.

That executive should sit in marketing. Not because the work belongs to marketing — it plainly spans product, data, and technology — but because marketing owns the outcome the work

produces. How a brand is discovered, described, and compared has always been marketing's result to answer for, and the agentic web changes the mechanism rather than the accountability. The functions that control the inputs stay responsible for them. Marketing is answerable for what the customer meets.

The test of the model is what happens when something goes wrong. Picture a campaign that describes the brand inaccurately. The creative lead may have made the mistake, but the executive is accountable if nobody supplied clear brand guidance, a reliable source of product facts, or a review step. The organization's job is to trace the failure to the broken input, whether that's unclear guidance, inaccurate content, outdated product information, conflicting data, or a missing review, rather than to the nearest person.

That traceable view of who owns what is an accountability map ↗: a shared record of who owns brand guidance, product truth, data access, and correction when a signal fails. It is the difference between a group that shares the work and a committee that dissolves the responsibility.

The Clarity Framework:

Ownership

THE QUESTION

Who is accountable for keeping the system aligned? Start by listing everyone who touches part of this work, across brand, content, product, data, technology, and communications, then narrow to one accountable executive and the small group of named owners answerable for their own inputs.

Who sets the brand rules the organization follows?

Who makes sure those rules reach creative, content, product, and data?

Who maintains the canonical source those teams draw from?

Who checks that published work still holds up once it appears in AI answers?

When something goes wrong, can we trace it to a decision, a missing rule, or a process failure, and does the person who owns that input know they own it?

Who decides when the answers to the other four questions need revisiting?

The Risk Of Deferral

Unassigned work doesn’t stop; it gets done inconsistently. Autonomy boundaries go unwritten, measurement runs without a baseline, the canonical source drifts, and none of it surfaces as an ownership problem. It surfaces as four unrelated operational problems, a year apart, with no one positioned to connect them.

The Operating Standard

One marketing executive accountable to the C-suite for the brand outcome, supported by named leaders in content, product, and data, each answerable when a failure traces to their area. Decision rights are explicit, a shared source of truth exists, and there is a documented path from a failure to the input that caused it. The group shares the work; accountability doesn’t disappear into it.

The Review Cadence

Revisit the map whenever the organization restructures, adopts a new class of agent, or changes what the systems can reach, and at least annually against the other four questions. If any of the four has no named owner, the map is out of date.

The Decisions A Competitor Can't License

Ninety-five percent of respondents report familiarity with generative engine optimization, ninety-three percent with answer engine optimization, and eighty-seven percent with Model Context Protocol. Sixty-three percent say they are aggressively integrating and scaling AI. Only two percent have no agents working anywhere in the function.

What hasn't kept pace is the operating structure underneath the terminology. Marketing owns brand visibility in AI systems at forty-five percent of organizations and nobody owns it at six percent. New AI titles are under consideration everywhere while decision rights go unassigned. Teams measure without a baseline, and grant autonomy without a written boundary.

CONCLUSION

“One of the ways that we use these agents is to better talk to humans, and to talk to other agents... The task of the marketer has always been, how do we reach our buyer, our consumer? How do we think about the person that we’re trying to talk to?”

— GABRIEL DILLON, GO-TO-MARKET LEAD, PERSONALIZATION, CONTENTFUL

Some of that is structural rather than careless.

Marketing is routinely required to take a public position on a new technology before anyone has had time to become expert in it, and this research caught the function mid-adjustment. The cost is that problems are identified late, and in an environment changing this quickly, late is expensive.

A competitor can license the same models. It cannot license a set of decisions an organization has already made, written down, and learned to operate.

So ask the five questions. Don’t grade the answers by confidence. This research has no shortage of it. Grade them by names, policies, baselines, sources, and proof. Then put them back on the calendar, because the systems they concern will not hold still, and a framework used once is just a meeting. ●

How We Did This Research



Marketing Decision-Makers

QUANTITATIVE.

350 marketing decision-makers surveyed across the United States (n=200), the United Kingdom (n=100) and Australia (n=50), fielded Q2’26 by Cint. Ninety-three percent are director-level or above (n=327); twenty-six percent are C-suite. Eighty-five percent work in marketing departments of more than fifty people (n=297); fifteen percent in departments of fifty or fewer (n=53). Forty percent are at companies with \$500M or more in annual revenue; thirty-seven percent directly control marketing budgets of \$1M or more.

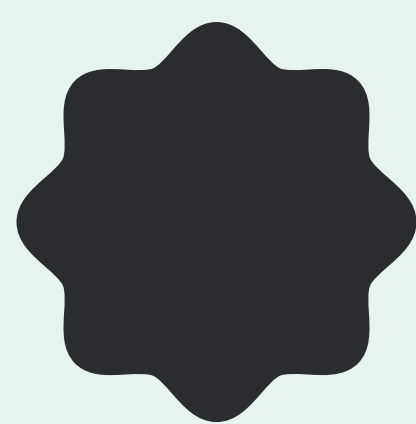
Respondents who said their organizations were not prioritizing AI were screened out of the survey. Our report aims to understand what happens after the investment decision, not whether it has been made. The tradeoff is that these findings describe organizations already committed to AI rather than the marketing profession as a whole, and figures throughout should be read against that base.

Differences between small and large marketing departments are reported as statistically significant only where the source tables mark them so.

QUALITATIVE.

A live, text-based conversational session conducted on the Remesh platform with 22 marketing executives independent of the Cint survey sample, supplemented by subject-matter interviews with Contentful executives. Qualitative figures are drawn from that group of 22 and identified as such wherever they appear. They are directional by design: Their value is in what they reveal about how leaders describe their problems when they aren’t choosing from a list.

These terms are the conditions currently observable in marketing organizations, with supporting evidence.



The Approval Tax

The time and value lost when an agent capable of routine work waits on a review from a human. Not all reviews are a tax — a review that catches a wrong claim before a customer sees it is the cheapest thing the organization does. The tax is what accrues on the approvals that exist because no one drew a line. Twenty-one percent of organizations require human approval before any agent action executes.



Decision Boundary

The written line between work an agent can do, work a person reviews, and work a person must decide. Forty-three percent of teams would let agents run across most tasks with periodic review, thirty-four percent confine them to low-stakes decisions, and twenty-one percent approve everything first. Those three groups are not disagreeing about agents. They are disagreeing about risk, mostly without having said so out loud.



Judgement Debt

The future cost of automating the entry-level work through which senior judgment has always been developed. Twenty-five percent of organizations have paused or reduced entry-level hiring and twenty-one percent have eliminated roles outright. The invoice arrives in about five years, addressed to whoever is running the department by then.



Canonical Brand Source

The one governed place where a company’s facts, claims, product details, and approved language live, which every downstream system is required to draw from. Ninety-five percent agree codified identity will widen a brand’s lead. Thirty-four percent have prioritized making brand information consistent across third-party sources.



Representation Drift

The slow change in how AI systems describe, compare, or recommend a brand, visible only against a prior record. Fifty-five percent of teams measure how they are described; thirty-two percent have checked once or twice, which records a single day and cannot show change at all.



Signal Integrity

Whether a brand’s facts, claims, and proof hold up across every source an AI system can reach. Thirty-seven percent of executives name the structure and clarity of content as what distinguishes a brand once AI is summarizing it; twenty percent name third-party validation. Signal integrity is the condition, and the canonical brand source is what produces it.



Accountability Map

A shared record of who owns brand guidance, product truth, data access, and correction when a signal fails. Marketing holds primary responsibility for AI visibility at forty-five percent of organizations; a dedicated AI, data, or innovation team at twenty-one percent; six percent have no clear owner. The map is the difference between a group that shares the work and a committee that dissolves the responsibility.

Atlantic Re:think ×  contentful