

Q4 2025 · OCTOBER — DECEMBER

The Credibility Problem

*Why Good Research Keeps Failing to Land — and What the Quarter's Nine
Conversations Reveal About the Gap*

performed signal

A Shared Diagnosis

Across the quarter's nine conversations, a single through-line emerged that no individual episode had been commissioned to explore: the gap between research quality and research impact has been misfiled as a methodology problem. Its real address is trust, organizational nerve, and what a business will let itself act on — and the industry has spent its effort on the wrong half. Better samples, faster turnaround, AI-augmented synthesis: these answer questions the field is no longer failing. The failures are upstream — Does the business trust the function? Can a researcher earn a decision, not just describe one? Is the right kind of problem even being researched?

The quarter offered no single fix — only a diagnosis its nine conversations kept independently arriving at.

Through-lines this quarter:

- **Trust precedes methodology.** Organizational fear — specifically, accountability anxiety — is the gate through which research must pass before any finding can influence a decision. The gate cannot be unlocked by methodological rigor alone.
- **Calibration is a lost discipline.** The digitization of research removed embedded quality controls that pre-digital fieldwork carried automatically. The industry has not replaced them; it has learned to live without them.
- **Analog is resurging, not retreating.** The dominant digital-transformation narrative has it backwards: digital infrastructure is scaffolding for analog experience, not a replacement for it. Consumer behavior's deepest drivers — convenience, physical presence, embodied preference — remain stubbornly analog.
- **Clarity is the delivery mechanism.** Data abundance is creating a clarity deficit, not a knowledge advantage. The researchers who land their work are the ones who have done the reduction before walking in the room.

- **AI's ceiling is real and locatable.** Synthetic respondents used as a proxy for human panels produce the outputs of research without the credibility that makes them worth acting on. The distinction between AI that synthesizes existing knowledge and AI that simulates novel human judgment is the whole game — the line between useful augmentation and what one contributor called “research theater.”

Five Structural Findings from the Quarter's Corpus

I. TRUST IS UPSTREAM OF METHODOLOGY

The most consistent finding across the quarter is structural: organizational trust is the gate through which research must pass before any finding can influence a decision. Methodological excellence does not unlock that gate; relationship, relevance, and a clear-eyed understanding of the organizational psychology on the other side of the table do.

From inside a fast-moving game studio, the diagnosis runs to a precise and uncomfortable place: fear.¹ Stakeholders who receive research findings face an accountability transfer. If they act on the data and the decision fails, ownership of the failure has shifted. The rational response to that risk — for a stakeholder whose political capital is limited — is to receive the findings without acting on them. The research was consulted; nothing was decided on it. The researcher becomes the organization's alibi, not its compass.

This is not distrust as professional skepticism. It is distrust as accountability anxiety. And it explains a pattern that many researchers experience without being able to name: meticulous work that produces a polite nod and then disappears. The problem is not the work. The problem is the transfer of risk that action would require.

The organizational politics angle sharpens the diagnosis: research quality is necessary but not sufficient.² The bottleneck sits one level past the findings' validity, in the organization's collective capacity to act on them. Global clients with distributed decision-making structures face a coordination problem that no research methodology can solve. The empirical question gets answered; the political question determines whether anyone acts on the answer.

Both vantage points triangulate toward the same prescription: trust-building is the real work, not soft preliminary work that precedes it. A researcher who has not built the organizational relationships to get findings acted on has not finished the

research. The deliverable is a decision, not a report.

The practical architecture of trust starts earlier than most researchers intervene. A Senior Research Manager at Electronic Arts argues that the work before the work is getting the problem type right.¹ She distinguishes between puzzles — problems that have correct answers, even when those answers are hard to find — and mysteries, which are structurally ambiguous and may not resolve. The industry's instinct is to treat every brief like a puzzle: commission a study, expect a clean answer. When a mystery is researched like a puzzle, the result is ambiguity where certainty was promised, and the trust cost of that failure compounds.

The corollary is tactical but non-trivial: research instruments that try to answer every adjacent question in a single study usually answer none of them well. Scoping to what a single study can credibly deliver is itself an act of trust-building — it trades the appearance of comprehensiveness for the reality of precision.

II. THE CALIBRATION GAP — WHAT DIGITIZATION REMOVED AND DID NOT REPLACE

If the trust problem is about the organizational conditions for research impact, the calibration problem is about the methodological conditions for research reliability. They are related — uncalibrated research erodes trust — but they originate in different failure modes.

A behavioral science firm working at the intersection of packaging research and shopper behavior puts the methodological principle plainly: what people do does not align with what they say, or what they say they do.² The response is to rely more heavily on observation — eye-tracking, physiological response, behavioral trace — than on stated preference. Three of the four stages in a framework for measuring a shelf encounter's effectiveness are behavioral rather than attitudinal. The reason is epistemological: behavioral data has a lower contamination rate from social desirability, question framing, and recall error.

But this principled commitment to behavioral observation does not, by itself, address the deeper calibration problem that the quarter's most historically grounded voice raises: beyond being observational, the auditing infrastructure of pre-digital research was structural.³ A field researcher administering a survey in person was simultaneously collecting the data and auditing the respondent's engagement. The fraud resistance was embedded in the method, not added to it afterward.

Digitization removed that friction. It also removed the quality controls that friction carried. Online panel research is cheaper, faster, and more scalable. It is also vulnerable to bots, duplicate respondents, satisficers, and speeders in ways that clipboard-era research was not. The problem is not that these vulnerabilities exist — the industry knows they exist — but that the auditing infrastructure has not kept pace with the attack surface.

The prescription that emerges is counterintuitive: being consistently wrong in a known direction is more useful than being accurately wrong in an unknown direction.³ A forecast model with consistent methodology produces errors that can be identified, corrected for, and priced into decisions. A forecast model that chases the newest panel source, the cleanest-looking data, and the most recent sample design produces errors that are structurally unpredictable. Far from being a compromise, consistency is the foundation on which calibration becomes possible.

This argument extends to AI tooling. Sentiment analysis at scale can distinguish positive from negative with reasonable accuracy.³ It collapses the dimensionality of emotional response in ways that matter: anger reads differently from disappointment; resignation reads differently from contempt. A Microsoft research director draws the capability line more precisely: AI is genuinely useful for synthesizing existing data, surfacing patterns in large qualitative archives, and accelerating the work that does not require novel human engagement.⁴ The decisive distinction is qualitative, separating AI that respects the epistemological stakes of a research question from AI that obscures them — the AI-versus-no-AI framing misses it entirely.

The calibration argument and the AI-ceiling argument are pointing at the same structural gap: the industry's quality infrastructure has not kept pace with its data-collection infrastructure. More data, faster, does not solve for the calibration that gives that data decision-weight.

III. THE ANALOG RESURGENCE — WHAT DIGITAL DISPLACED AND CANNOT REPLACE

Kantar's strategy lead opened his Q4 episode with a claim designed to stop the room: the future of digital is analog.⁵ By analog he meant the terminal destination of the digital build-out, with no managed coexistence implied.

The argument is not nostalgic. It follows a logic about infrastructure and experience. Digital systems are built to serve human preferences, and human preferences are, at their deepest register, embodied, social, and physical. The digital layer is scaffolding; the analog experience is what the scaffolding was built to support. The conclusion — that digital saturation will eventually drive a revaluation of analog experience — is a prediction about where consumer desire will settle once the novelty of digital access normalizes.

The evidence the quarter's conversations surface is inconvenient for technology optimists: despite decades of digital saturation, convenience remains the single largest driver of consumer choice.⁵ Convenience is not a digital property. Amazon is convenient. So is a corner bodega you have been going to for twenty years. The digital and analog versions of convenience serve the same underlying preference — friction reduction, reliability, fit with a life already in motion. Technology is instrumental to the experience; the experience is what the consumer is actually purchasing.

This reframe has direct methodological consequences. If the most durable consumer preferences are analog in character, then research frameworks built on digital interaction data will systematically underweight the experiences that drive the most consequential purchase decisions. The shelf encounter — physical, sensory, momentary — is still the decisive consumer moment in food retail, as the quarter's

food industry contributor confirms.⁹ The category-specific research apparatus built to understand that moment is a stress test for the broader thesis: analog drivers deserve more research investment than a purely digital framework prescribes.

The resistance to change-obsession that Kantar's strategy lead articulates is a methodological conservatism, not a technological one.⁵ His precise claim is that AI is helping consumers do things in the same ways, faster — not inventing new things for them to do. The shopper journey looks like the shopper journey. AI search and AI recommendations operate on a fundamentally unchanged set of consumer objectives. The research question remains the same; what changes is the data source and the synthesis tool. Understanding what consumers want, why they want it, and how a product fits into their life remains an unsolved, un-obsolete problem — one AI has only made faster to work on.

The implication for the research profession is about where durable value lives. If AI absorbs the execution layer — data collection, synthesis, preliminary pattern recognition — what remains is the interpretive judgment that determines whether a pattern is meaningful, whether a finding is actionable, and what the question should have been in the first place. The image one contributor returns to is deliberate: Rodin's Thinker, bearing full weight on the act of sustained intellectual attention.⁵ Not faster thinking, not better-tooled thinking. Just thinking.

IV. THE CLARITY CRISIS — WHY MORE DATA MAKES THE PROBLEM WORSE

The industry's instinctive response to the trust deficit is more evidence. More data, more research rounds, more confidence intervals. The quarter's sharpest communications-side diagnosis is that this response misreads the problem at a structural level.⁶

CMOs suffer a clarity shortage, while data is the one thing they have in surplus. The distinction matters because solving a clarity problem with more information compounds the deficit. Every new report, every additional dashboard, every

expanded research cadence adds to the cognitive load on the executive who must translate all of it into a decision. The research function, trying to demonstrate its value through comprehensiveness, is often doing the opposite.

The acid test that a founder of a CMO advisory practice proposes is sharp: if you cannot distill a finding to one slide, you do not yet own the story.⁶ This is not a presentation heuristic. It is an epistemological claim about comprehension. The ability to collapse a finding to its essential structure is evidence that you have understood it at depth. The researcher who needs seven slides to convey a finding probably needs another week with the data.

The same discipline arrives independently from a strategy and consulting perspective at Kantar.⁵ The best presentations are about one thing. Every slide makes one point. The convergence — reached from different seats in different sectors by different routes — is Q4's highest-confidence echo. Two independent vantage points landed on identical formulations of the same principle without coordination. When the indexed record surfaces that kind of independent convergence, the claim carries more than editorial weight.

Inside a large financial services organization, the clarity argument takes a different form: the executive attention economy.⁷ Researchers who get heard are the ones who have done the reduction work before walking in the room. The failure mode she identifies is the completeness instinct — the researcher's urge to present every caveat, every nuance, every hedged finding. Comprehensiveness signals rigor to another researcher. To an executive deciding between competing demands on a Tuesday afternoon, it signals that the researcher has not decided what matters.

The seat-at-the-table argument that a CX research principal at Mutual of America makes is structurally related but arrives via a different diagnosis.⁸ His prescription is curricular rather than communicational: understand the business deeply enough that research is anchored in its actual risk architecture. Far from contradicting the CMO advisory founder, he completes the picture — together they map the complete competency the quarter demands. A researcher who understands the business deeply and can speak to it precisely. The first enables the second; without both, neither lands.

The food industry variant of this credibility problem operates at the intersection of high SKU counts, fragmented consumer segments, complex private-label dynamics, and a health-claims landscape that requires evidence quality the broader insights industry does not always deliver.⁹ The organizational trust framework that the quarter maps is not sector-specific. It lands with equal force in CPG as in gaming, financial services, and behavioral science.

V. THE SYNTHETIC DATA QUESTION — WHERE THE CAPABILITY CEILING IS

The quarter's sharpest challenge to AI-augmentation optimism comes from inside a major technology company, where a director of market research draws a capability line that the broader field has been reluctant to draw precisely.⁴

The practice he targets is specific and growing: using synthetic respondents — AI models trained on prior human responses — as a proxy for real human panels. The output of this practice produces the artifacts of genuine research: charts, cross-tabulations, significance flags, and executive-ready topline. It does not produce what those artifacts are supposed to represent: the novel responses of actual humans to questions they have not previously been asked.

The phrase he coined for this practice — “*research theater*” — names something the field has been circling without quite landing on. Research theater is the performance of research without its epistemological substance. It is the appearance of insight without the act of insight. A model interpolating from prior human responses synthesizes what humans have already said rather than simulating what they would say — a different and lesser epistemic operation, particularly when the research question involves novelty, category disruption, or experiences the training data has not indexed.

The accountability stakes sharpen the concern. When research is used to justify a decision and that decision fails, credibility is on the line — and so is the research basis for the decision. Synthetic respondents produce evidence that cannot bear that

weight, not because the interpolations are necessarily wrong, but because the chain of epistemic custody that validates evidence in a high-stakes business decision is broken.

The broader argument for AI is not anti-augmentation. AI that synthesizes existing qualitative archives, surfaces patterns in large datasets, and accelerates preparatory work that does not require novel human engagement is genuinely useful. ⁴ The line is between AI as a tool that extends what researchers can know from what has already been collected, and AI as a replacement for the collection of new human judgment. The first is augmentation. The second is substitution — and when the research question requires novelty, substitution is a defect, not an efficiency.

The prescriptive framing that emerges from this vantage point is about the relationship between the researcher and the tool: AI as sparring partner rather than butler. ⁴ A butler executes without challenging. A sparring partner stress-tests. The research relationship with AI should be adversarial by design — the tool pushes back on the researcher's assumptions rather than validating them. Read correctly, that inversion is a disposition about what AI is for in a research workflow, well beyond any prompt-engineering trick.

Where the Field Contradicted Itself

The quarter's convergences are notable. So are the places where the field contradicted itself — where equally grounded vantage points reached different conclusions from similar evidence.

Speed and stability cannot both be maximized. The research function faces competing demands that the quarter does not resolve: a Senior Research Manager at a major game studio argues that research operating at anything less than business-decision speed becomes irrelevant by definition — by the time it lands, the decision has already been made. ¹ A veteran forecaster with decades of calibration experience argues, with equal force, that the pursuit of speed is trading consistency for the appearance of responsiveness. ³ Calibration requires a stable methodology across time; instability — even in service of accuracy — destroys the error-prediction infrastructure that makes forecasts trustworthy. Both arguments are correct within their domains. The tension between them does not resolve at the level of abstraction at which it is being argued. It resolves in the specific decision being researched, the specific organization commissioning it, and the specific cost of being wrong in each direction.

Where AI's useful territory ends is not yet agreed. A behavioral science executive argues that human validation at the final executive decision stage will remain mandatory indefinitely: the C-Suite needs a chain of epistemic custody that traces back to real consumers. ² Kantar's strategy lead holds that AI is currently doing existing things faster rather than inventing new things, but does not foreclose the possibility of future categorical shifts. ⁵ A Microsoft research director draws the ceiling more precisely at the boundary between synthesizing prior data and simulating novel human judgment. ⁴ These three framings do not directly contradict one another, but they triangulate differently on how far AI's useful territory extends — and whether that boundary is permanent or provisional.

Whether fundamentals or transformation is the dominant story is unresolved. The quarter's most pointed methodological conservatism — the claim that the core job of the insights function has not changed, only the tools — stands in productive tension with the field's prevailing transformation narrative.⁵ The conservative case is that durable strategic value lies in the things that don't change: consumer preferences, organizational dynamics, the irreducible need for interpretive judgment. The transformation case is that the pace of AI development will eventually cross thresholds that make existing research infrastructure obsolete in ways the quarter's conversations have not yet anticipated. The indexed record is tracking this as an open empirical question.

Vocabulary That Gained Currency This Quarter

Research theater

“research theater”

A research practice that produces the outputs of genuine inquiry — charts, cross-tabulations, significance flags — without its epistemological substance. Specifically applied to synthetic respondents as a proxy for human panels: the interpolation of prior responses is a different and lesser epistemic operation than eliciting novel human judgment.

Microsoft research director ⁴

The puzzle/mystery distinction

“A puzzle is relatively easy to solve... A mystery is complicated, ambiguous”

A problem-type taxonomy for research design. Puzzles have correct answers, even when those answers are hard to find; mysteries are structurally ambiguous and may not resolve. Treating a mystery like a puzzle — commissioning a study and expecting a clean answer — produces the characteristic failure mode of findings that disappoint the stakeholder and erode trust.

Senior Research Manager, Electronic Arts ¹

Research as trust architecture

The framing that organizational trust is not background context for research but its primary delivery mechanism. A researcher who produces methodologically sound findings that the business does not trust has not completed the research; the deliverable is a decision, not a report.

Synthesized across Electronic Arts ¹ and Behaviorally ² vantage points

AI as sparring partner

“AI as a sparring partner, not a butler”

A prescriptive relationship model for researcher–AI interaction. A butler executes instructions without challenging them; a sparring partner stress-tests assumptions. The distinction implies that AI’s most valuable role in a research workflow is adversarial — surfacing counterevidence and pushing back — rather than validating or automating.

Microsoft research director ⁴

The golden age of analog

“We are about to enter a golden age of analog”

The thesis that digital infrastructure is scaffolding for analog experience rather than a replacement for it — and that digital saturation will eventually drive a revaluation of embodied, physical, and social experience. Applied to the research profession: analog consumer drivers (convenience, physical presence, embodied preference) deserve research investment that a purely digital framework systematically underweights.

Kantar Knowledge Lead for Strategy and Consulting ⁵

Calibration discipline

“It’s okay to be wrong as long as you’re consistently wrong”

The methodological practice of maintaining stable research procedures across time so that errors become predictable and correctable. Distinguished from accuracy: a model that is consistently wrong in a known direction is more useful for decision-making than one that varies its methodology in pursuit of accuracy, because the latter produces errors that cannot be priced into decisions.

Veteran product launch forecaster, CG Research & Consulting ³

The clarity gap

The structural deficit produced when data abundance outpaces the research function's capacity to reduce findings to their essential stakes. Specifically: CMOs are clarity-starved, not data-starved, and adding more research to the problem compounds the deficit rather than resolving it.

CMO advisory founder ⁶

Accelerating, Subsiding, and the Year Ahead

ACCELERATING

- ▲ **The behavioral observation premium.** Stated-preference research is being progressively displaced in high-stakes decisions by behavioral and observational methods — eye-tracking, physiological response, purchase trace data. The convergence on “what people do, not what they say” across multiple sectors in Q4 suggests this is becoming a baseline expectation rather than a methodological differentiator.²
- ▲ **The clarity mandate.** As AI accelerates data synthesis and report generation, the value of reduction — the ability to collapse a complex finding to a single, falsifiable, decision-relevant claim — is rising. The researchers who advance are those who have mastered reduction before walking in the room.^{5 6}
- ▲ **Role-and-org attribution as analytical currency.** The quarter’s strongest arguments were inseparable from the institutional vantage points that produced them. A Microsoft research director’s synthetic-data skepticism carries different weight than a generalist commentator’s. Kantar’s strategy lead’s analog thesis lands differently than a challenger firm’s would. The seat is part of the signal.

SUBSIDING

- ▼ **The comprehensive deliverable as a proof of value.** The research function’s traditional defense of its value — thoroughness, coverage, hedged findings, executive-ready appendices — is losing credibility with the executives who must act on research. The expectation of reduction is displacing the expectation of completeness.^{6 7}
- ▼ **Uncritical AI optimism.** The quarter registered a measurable shift from “AI is transforming research” to “AI has a locatable ceiling and we should name it.” The coinage of “research theater” as a critique of a specific synthetic-data practice signals that the field is developing vocabulary for what AI should not be used for — which is a different epistemic register than enthusiasm.⁴

- ▼ **Methodology—first framing of the researcher’s value proposition.** The consistent throughline across Q4 is that methodological rigor, while necessary, is not the primary determinant of research impact. The profession is recalibrating — slowly and unevenly — toward relationship, organizational literacy, and communication as the primary competencies that enable findings to land.

THE YEAR AHEAD

Synthetic respondent practices will face an institutional reckoning, not a market correction. Usage of AI-simulated panels will continue to expand in 2026, but expect at least one significant, publicly disclosed research decision — a product launch, brand pivot, or communications strategy — to trace a costly failure back to synthetic rather than real respondents. When that failure is attributable, the conversation will shift from epistemological to liability, and procurement standards will begin to catch up.

Signal: a Microsoft research director’s “research theater” framing was adopted by peers in Q4 without attribution — the critique is diffusing beyond its origin point. ⁴

The one-thing discipline will enter practitioner training curricula. The convergence of two independent vantage points on the single-slide / single-point principle — reached without coordination, from different sectors, framed in nearly identical terms — suggests this is approaching consensus rather than preference. Expect to see it articulated as a competency standard in at least one major insights-industry body’s practitioner framework by end of 2026.

Signal: independent convergence from a CMO advisory founder ⁶ and Kantar’s strategy lead ⁵ in the same quarter, using nearly identical formulations.

Analog experience research will receive dedicated investment lines in at least one major consumer goods category. If the analog resurgence thesis is correct, the research apparatus in categories where physical experience is the decisive consumer moment — food retail, personal care, apparel — will require methods that digital interaction data cannot provide. Expect budget reallocation toward observational and ethnographic methods in at least one major CPG firm's 2026 research planning cycle, cited explicitly against digital over-indexing.

Signal: the analog thesis surfaced from Kantar's strategy lead ⁵ and was independently confirmed by the food industry vantage point ⁹ — two separate sector perspectives landing in the same place.

Trust-building will be formalized as a research competency, not treated as a soft skill. The quarter's convergence on organizational trust as the primary gate for research impact — reached independently from a gaming studio, a behavioral science firm, and a financial services UX operation — suggests the profession is ready to codify what it currently treats as interpersonal luck. Expect at least one major insights association to add a trust-architecture track or session stream to its 2026 conference programming.

Signal: the trust-as-gate argument surfaced from three independent institutional vantage points — Electronic Arts ¹, Behaviorally ², and Prudential Financial ⁷ — without coordination.

Q4 2025 Contributors

- [1] **Leena Doshi**, Senior Manager of Research, Electronic Arts — Trust architecture, puzzle/mystery framework, survey scope discipline (Q4 2025). *[Episode link placeholder]*

- [2] **Matthew Salem**, SVP of Customer Success, Behaviorally — Behavioral observation, packaging research, organizational politics in decision-making (Q4 2025). *[Episode link placeholder]*

- [3] **Charlie Grossman**, Owner, CG Research & Consulting — Calibration discipline, pre-digital quality infrastructure, AI sentiment limits (Q4 2025). *[Episode link placeholder]*

- [4] **David Evans**, Director of Market Research, Microsoft — Synthetic data skepticism, research theater, AI as sparring partner (Q4 2025). *[Episode link placeholder]*

- [5] **J. Walker Smith**, Knowledge Lead for Strategy and Consulting, Kantar — Analog resurgence thesis, change-obsession critique, The Thinker, one-thing principle (Q4 2025). *[Episode link placeholder]*

- [6] **Steve Olenski**, Founder, CMO Whisperer Advisory — Clarity deficit, single-slide discipline, insights as growth engine (Q4 2025). *[Episode link placeholder]*

- [7] **Amberly Miller**, Director of UX Research, Prudential Financial — Selectivity in executive communication, completeness instinct failure mode, AI augmentation ceiling (Q4 2025). *[Episode link placeholder]*

- [8] **Ed Kahn**, Senior Principal of CX Research and Insights, Mutual of America — Business literacy as research competency, seat at the table, CX research ROI (Q4 2025). *[Episode link placeholder]*

- [9] **Steve Markenson**, Vice President of Research and Insights, FMI — the Food Industry Association — Food retail insights, analog shelf encounter, AI in CPG context (Q4 2025). *[Episode link placeholder]*

Notes & Methodology

Indexed record as instrument. This white paper treats the indexed record of the quarter's nine episodes (the *corpus*) — a structured set of professional perspectives analyzed for convergence, tension, and emergent vocabulary — as the instrument of analysis, not the subject. The analysis surfaces what the indexed record reveals collectively that no individual episode could reveal alone.

Attribution method. All findings are attributed by role and organization rather than personal name in body prose, consistent with the indexed-record-as-instrument frame: the analytical weight of a claim derives from the institutional vantage point it was produced from, not from the individual. Full names appear only in the References & Contributors section. Organization affiliations reflect each contributor's role at the time of the episode.

Source fidelity. No claim, statistic, or quoted phrase in this white paper is introduced from outside the existing verified episode transcripts. All verbatim coinages in the Emergent Lexicon section are drawn from reviewed transcripts. This is a restructuring and reanalysis of existing material, not a re-research.

Cross-reference basis. The convergences and tensions identified in this paper are drawn from the Curiosity Current indexed episode record, which maps cross-speaker claim relationships — echoes, agreements, tensions, extensions — across indexed episodes. Quarter-scoped counts (e.g., “the quarter's nine conversations”) refer to a closed historical set; no living totals are asserted.

Episode links. Full episodes are available on the Curiosity Current podcast feed on Apple Podcasts, Spotify, and YouTube.

Curiosity Current is produced by AYTM. This white paper was compiled by the Curiosity Current editorial team from the quarter's indexed episodes. Published Q1 2026.

Curiosity Current · AYTM · Q4 2025



© 2026 aytm. All rights reserved.