

Q1 2026 · JANUARY – MARCH

The Synthesis Problem

When Cheap Outputs Undermine Costly Truth

The Synthesis Problem: When Cheap Outputs Undermine Costly Truth

A research procurement team approves a synthetic–respondent platform. The demo shows fluent transcripts, clean cross–tabs, and a 95% correlation against a benchmark survey. The platform clears procurement, the brand ships against its findings, and six weeks later the launch underperforms. When the retro asks the obvious question, was the data wrong, there’s no clean answer to give. The respondents may have existed, or they may not have. The 95% may have meant the synthetic captured reality, or it may have meant the synthetic and the benchmark were wrong in the same way. The chain of epistemic custody is broken, and no one in the room knew when it broke.

Practitioners kept returning to that moment across the quarter’s eleven Q1 2026 conversations. From synthetic–data platforms, global media measurement, food–culture consultancy, academic food science, corporate insurance, and healthcare strategy, we heard a single worry in different accents: the tools that make research faster and cheaper are, in the same motion, eroding the epistemological legitimacy of what they produce. Speed gets mistaken for insight, plausibility for accuracy, procedural correctness for a trustworthy conclusion. The quarter’s worry isn’t AI itself but the prior question it keeps forcing: what now makes a research claim real? Q1 returned an answer in eleven independent voices and one converging shape, five gaps the field has been treating as one, each requiring a different repair.

Through–lines:

- Plausible versus checkable. When LLM–generated outputs look like research outputs, the only thing that separates them is an audit the digitization of research quietly removed. The synthetic data debate has reached a definitional breaking point, and procurement is evaluating categorically distinct capabilities, causal scenario modeling versus fabricated respondents, as if they were equivalents.

- **Statistic versus story.** A finding lands in the decision room when the listener arrives at the conclusion before it's stated. Research delivery that climbs a pyramid of evidence loses every decision-maker who reads only the top of it.
- **A number versus a replicable number.** Methodological guardianship has emerged as an explicit professional identity claim: researchers framing themselves not merely as analysts but as defenders of replicability against tool-mediated erosion.
- **Reluctant versus selective.** Consumer segmentation frameworks are failing at their edges, both for older adults and for ideologically disparate groups converging on shared purchasing behaviors. The frameworks aren't getting smarter and the consumers haven't changed; the instruments have stopped surfacing the surprise.
- **Study completed versus decision influenced.** The organizational distance between research production and decision influence remains the industry's most persistent unsolved problem, and the practitioners who've closed it describe a common prerequisite: clarity about what a researcher brings that AI doesn't.

Five Convergent Shifts Across Q1 2026

I. Plausible vs. checkable

The most famous quality failure in consumer research history is also Q1's clearest illustration of the gap. The New Coke decision tested cleanly in quantitative aggregation, with net preference breaking favorably and focus-group dissent treated as statistical noise. It shipped into one of the great category retreats of the century. The methodology was procedurally correct; the interpretation was catastrophically wrong. Procedural correctness, the case shows, is the failure mode the field has been mistaking for the solution. A food marketing academic surfaced the case in Q1 to make a specific argument: pre-register hypotheses, split datasets, treat qualitative dissent as signal rather than noise.⁸ The prescription is a guardianship argument applied to study design, and it lands four decades later because the gap has only widened. Today's version of the New Coke decision is built faster, costs less, and breaks the same way.

The synthetic data debate reached a definitional breaking point in Q1. The founder of a synthetic-data platform spent much of his episode less advocating for his category than distinguishing it from what most of the market means when it uses the term.¹ The conflation has consequences. "Synthetic data" now covers both causal AI that models relationships within genuine survey datasets and LLM-generated respondents that interpolate from training corpora, two capabilities with radically different epistemological legitimacy, evaluated in procurement as if they were the same thing.

The epistemological problem with LLM-generated synthetic respondents is specific and structural. These systems produce outputs that reflect the distribution of text that existed before the survey was designed, so they can't surface genuinely novel consumer positions. And they fail in ways that are difficult to detect from the output alone: the same question posed in English versus Spanish activates different parts of the network and returns materially different answers, with no signal in the result that anything has gone wrong.¹ The

synthetic-data founder drew a clarifying contrast to coding, the domain where LLMs have had their most commercially visible success. Wrong outputs in software fail visibly, get tested, and get retried. Consumer research has no equivalent sandbox. The output looks plausible whether it's accurate or not, because accuracy can't be checked against a ground truth the research was designed to discover.¹ An independent articulation of this failure mode appeared in Q4 2025 under the label "research theater," the production of outputs that perform the function of research without its epistemological content. Q1 has now located that failure at its methodological root.

The affirmative case for legitimate synthetic data is narrower and more defensible: causal scenario modeling applied to *existing real-respondent data* to extend analysis into subgroups the original sample was underpowered to reach. The respondents are real; the extended analysis is synthetic.¹ The distinction matters because it preserves the human signal at the base of the inference chain, the thing LLM-generated respondents lack by construction.

The underpowered study is the industry's most persistent and least discussed quality failure. The same synthetic-data platform founder cited a striking figure: more than 50% of surveys, he claims, return no statistically valid results. If that's accurate, the research industry's quality problem lies primarily in study design and powering, not in analysis or visualization.¹ This claim is flagged as unverified (see Notes & Methodology), but its directional force is consistent with what the quarter's other methodological voices observed from different angles: studies are routinely commissioned that can't, in principle, return actionable answers.

The "just-add-sugar" problem names the failure mode that follows. Regression finds that sweetness predicts preference, and it does. But the researcher is supposed to know that a product already at its formulation limit, or whose brand positioning requires restraint, can't act on that correlation. The regression can't make that judgment. When the researcher delegates interpretation to the algorithm, the category error becomes invisible.¹

The CEO of a leading panel quality firm extended the epistemological crisis to sampling, the layer of the research stack least visible to clients and most vulnerable to degradation.⁶ Bot responses, professional survey farmers, and AI-

assisted identity fabrication represent a systemic threat, not a nuisance. And the tools enabling faster research synthesis simultaneously lower the barrier to fabricating convincing-looking respondents. The industry is in an arms race it isn't winning.⁶ Clients who treat sampling as a procurement decision, awarded to whoever delivers fastest at lowest cost, create the conditions that fraud exploits.

Where this leaves us is a line procurement standards will have to start drawing in 2026. The legitimate tier models on real data and instruments AI to read the real signal more cleanly; the illegitimate tier inflates a small sample to half a million synthetic rows when the original survey already failed, and adds nothing the original couldn't deliver. The work is upstream of that. It takes a practiced human to tell apart a result that's true from one that merely looks plausible, and that skill is becoming the scarce, appreciating asset, not the depreciating one.

II. Statistic vs. story

A parent drives the family through a fast-food window on a hard day. A diet tracker logs the meal as failure. A research frame that captures context logs it as a successful act of caregiving under constraint. Same behavior, two readings. One produces a decision that lands; one produces a decision that doesn't.

The CEO of a food culture consultancy with decades of longitudinal consumer research brought the storytelling thesis to the say-do gap, market research's most persistent methodological challenge.¹¹ Her reframe is significant: rather than consumer deception, the say-do gap reflects a gap between aspiration and context-dependent action. If the gap is framed as deception, the research response is interrogation. If it's framed as aspiration interrupted by constraint, the research response is empathy, reading the situational pressures that interrupt expression rather than the stated preference alone.¹¹ Research that misses context misinterprets behavior, consistently, systematically, and in ways that generate real business errors.

Three Q1 practitioners approached the same question from different disciplines, what makes a finding land, and converged on a shared answer: narrative outperforms aggregate in the room where decisions are made, and the mechanism is arrival rather than persuasion.

The Corporate Vice President of Consumer Insights at a major life insurer made the claim with the most direct formulation: recounting a story someone shared in a meeting is more valuable than citing an aggregate percentage.² The point concerns persuasion mechanics, not the standing of statistics. The 82% is a fact; the story is an argument; decision-makers respond to arguments. The more counterintuitive finding is about how the story works. The most effective research delivery doesn't conclude with the insight at the bottom of a pyramid of evidence. It constructs a narrative that allows the listener to arrive at the conclusion before it's stated, so that when the conclusion appears, the listener already holds it.² Resistance collapses because there's nothing to resist. The listener hasn't been told what to believe; they've already believed it.

A Vice Provost for Innovation at a research university makes an adjacent argument through his methodological practice. Play-based research methods, specifically structured construction with abstract materials, reduce the social cost of genuine disclosure.³ Traditional qualitative research asks participants to speak, which activates self-presentation, social desirability, and the desire to give the right answer. Building sidesteps this: the hands are engaged, the conscious monitor is partially occupied, and the metaphoric distance of the build lets people express positions they might not verbalize directly.³ The deeper critique is about what research *for* means. Research that only reports what's true misses why it's true, how it became true, and what other truths it's connected to. Those adjacent dimensions are the ones that generate genuinely novel insight, the kind that changes how a problem is framed, not just which answer is selected.³

Empathy in this register isn't a values-deck word; it's a production input. It's the mechanism that determines whether a survey instrument captures aspiration interrupted by context or registers it as failure to comply. The instrument that reads the fast-food-parent moment as caregiving under constraint, not as a

lapse, is the instrument that produces the finding capable of landing in a real meeting. Apply empathy to the respondent the same way you apply it to the consumer behavior you're researching, and the story that moves the room writes itself.

III. A number vs. a replicable number

Q1 2026 is the quarter where a specific professional identity formulation appeared with enough frequency and independence to qualify as an emergent position: the researcher as guardian of methodological legitimacy against tool-mediated erosion.

Snap's global head of research and insights named the threat with particular clarity: the democratization of research tooling will produce practitioners who believe research is easy and don't appreciate the distinction between a number and a *replicable* number.⁷ Far from opposing AI in research, the same practitioner described active use of AI for knowledge management, addressing the field's chronic problem of insights that ship into organizational oblivion, never to surface again.⁷ The guardianship argument is precise: AI can execute against a well-formed question; it can't form the question. The act of deciding what you actually need to know, what level of specificity is required, what alternative framings have been considered and rejected, is itself a form of analysis. Delegating that act to a tool that optimizes for plausibility rather than precision is delegating the intellectual core of the work.⁷

The prescription follows from the diagnosis: the researcher's value in the AI era is the critical thinking that precedes questionnaire execution, rather than the execution itself, interpreting what findings mean for the business, holding organizational context, and advocating for conclusions in the room where decisions are actually made.⁷ And doing so without apology. The injunction that closed one Q1 episode, to be bold in professional voice, to stake a claim rather than retreat into hedged findings, reads as the operational corollary of the guardianship thesis: you can't defend the integrity of the work if you're equivocal about the value of the work.⁷

The founder of a research identity consultancy approached this from the identity angle, arguing that practitioners who'll thrive in the AI era are those who've clarified what they bring that AI doesn't.⁵ In a landscape where execution is increasingly automated, the differentiated researcher is defined by judgment: holding organizational context, navigating stakeholder dynamics, and advocating for findings where decisions are actually made. This functions as a prerequisite diagnostic, you can't stake a professional claim until you know what claim you're uniquely positioned to make.⁵

The guardianship theme has organizational implications beyond individual practitioners. A healthcare insights VP at a major professional association described the structural challenge of research navigation inside large institutions: findings must be made navigable within specific organizational architectures before they can influence decisions.⁹ The gap between having a finding and deploying it is where most research value is lost, and closing that gap is unglamorous operational work that neither methodological rigor nor compelling narrative alone can substitute for.⁹

Here's the contrarian reading the quarter points toward. The industry's default response to AI is efficiency-first: cheaper outputs, faster turnarounds, fewer hands. The guardianship claim cohering across Q1 argues for verification-first instead, with experienced humans as the scarce asset that appreciates rather than depreciates. As the ocean of plausible-looking insights rises, telling true from merely-true-looking takes a lifetime of practiced judgment, and the price of that judgment doesn't trend toward zero. It reverses its multi-year decline.

IV. Reluctant vs. selective

Adults 50 and over own smartphones at numbers near parity with the general population, a figure that's nearly doubled over the prior decade. The senior consumer insights manager at AARP arrived in Q1 with what reads at first as a small corrective and turns out to be a frame: what reads as reluctance among older adults is in fact selectivity.⁴ The distinction reframes the research question entirely. A reluctant consumer is a conversion problem; a selective consumer is a

design and positioning problem. The research interventions are completely different, and designing for the wrong problem guarantees a wrong answer. Ownership parity coexists with real adoption gaps, but the gaps survive on patterns for specific features and form factors, and on design that systematically fails to account for the accessibility and legibility needs of older users.⁴ The implication is a universal design argument with demographic evidence: features optimized for older users, larger text, clearer navigation, forgiving error states, benefit all users. The curb-cut principle applies to digital product design.⁴

The CEO of a food culture consultancy contributed the quarter's most consequential trend observation: three previously ideologically incompatible consumer cohorts are converging on shared purchasing behaviors around food quality.¹¹ MAHA (Make America Healthy Again) adherents, GLP-1 users reducing caloric intake, and progressive consumers focused on sustainability and sourcing have historically been opposed along political, dietary, and cultural lines. They're now arriving at overlapping conclusions through entirely different reasoning paths.

For food brands tracking volume as the primary metric, this convergence is invisible. For brands tracking quality dimensions of purchase decisions, it's a significant opportunity signal: consumers eating less may still pay more for higher-quality food, with volume losses offset by trading up.¹¹ The research design determines whether you can see it. The same practitioner offered the quarter's most counterintuitive recommendation for trend identification: befriend the dissatisfied.¹¹ The consumers most useful for surfacing what's coming are the dissatisfied minority, whose unmet needs still sit below the threshold of mainstream demand, rather than the satisfied majority. Research programs organized around optimizing for satisfied users are structurally blind to where the market is moving.

Two segmentation failures, two reframes, one underlying instrument problem. The corrective is methodological rather than tactical: the research design determines what you can see. A framework that lumps device ownership and feature adoption into "older adults are slow to adopt" describes the framework's blind spot, not the consumer's pace. A purchase-tracking framework that

aggregates to volume describes the framework's blind spot, not the consumer's wallet. The work shifts upstream of analysis, to the question of what the instrument was designed to surface in the first place.

V. Study completed vs. decision influenced

The organizational gap between research production and decision influence was a persistent preoccupation across Q1 2026, and the practitioners who've closed it describe a common set of conditions that made the crossing possible.

A Group Product Marketing Manager for Insights at a major professional network reduced the function's value proposition to a single test: every research question being answered should be traceable to a business decision it seeks to influence.¹⁰ Questions that can't be linked to a specific decision are, at best, calendar-filling. The test is diagnostic. It filters out work the organization doesn't actually need, and it gives researchers a frame for positioning their work that bypasses the perennial debate about research's "seat at the table." If the decision linkage is explicit, the seat is implied.¹⁰

The Corporate Vice President of Consumer Insights at a major life insurer articulated one of the quarter's most counterintuitive hiring arguments: category expertise can be a liability.² The researcher who's spent their career in a single industry already knows, or believes they know, what consumers in that category think. The certainty is itself a bias; familiarity makes it harder to notice what's surprising about the familiar. Researchers who arrive from adjacent categories ask the questions that insiders have stopped asking.² The same practitioner named a novel compound development challenge: COVID-era career interruption, accelerated by AI adoption, has produced a cohort of researchers with distinctive skill gaps in interpersonal consultation.² The professional socialization that builds networks, mentoring relationships, and the habit of collaborative inquiry was disrupted at the point of entry. These researchers may have stronger tool fluency and weaker interpersonal skills than any prior cohort, a gap the field hasn't fully named, let alone addressed.² The trait most predictive

of eventual influence, in this practitioner's experience, is intrinsic motivation to understand the business, the disposition to care what the numbers mean, not merely to produce them.²

The Vice President of Strategic Insights at a major medical association extended this to the specific complexity of healthcare stakeholders, where the decision-maker and the relevant audience are rarely identical.⁹ Physicians, patients, administrators, payers, and family members participate in healthcare decisions in relationships that standard segmentation frameworks aren't designed to hold simultaneously. The researcher who identifies the most visible person as the relevant decision-maker will design studies that miss the actual decision architecture.⁹

The through-line underneath these accounts is a discipline worth stating plainly: every research question should seek to influence a business decision. Studies built without a named decision in mind cost the same as studies built around one and return less. The most predictive trait isn't methodological depth but the disposition to ask what a finding would move. The work is upstream of the data, and the seat at the table follows.

Where Q1 Contradicts Itself

Causal modeling versus primary collection.

The synthetic-data platform founder and the panel quality CEO aren't as far apart as they first appear, both critique LLM-generated synthetic respondents, but their positions diverge on the substitutability of causal scenario modeling for primary data collection.^{1,6} The platform founder argues that causal modeling on real respondent data is categorically valid for specific research questions, particularly scenario extension into underpowered subgroups. The panel quality CEO argues that AI is “fit for purpose, not a silver bullet,” and her skepticism extends toward synthetic respondent use cases broadly. The degree to which properly constructed causal synthesis can substitute for primary data collection for defined questions remains the field's most consequential unresolved technical debate.

Where rigor risk lives.

The global insights head and the corporate insights VP agree that AI introduces methodological risk, but locate it differently.^{7,2} One frames the risk as professional: researchers may cede the intellectual core of their work, question formation and interpretation, to tools that optimize for plausibility. The other frames the risk as generational: the AI era has accelerated a development disruption in early-career researchers, replacing mentored consultation with tool-mediated lookup. Both are real. Neither is the complete account.

Consumer aspiration versus minority signal.

The food culture CEO's reframe of the say-do gap as aspiration interrupted by context sits in productive tension with the food marketing academic's New Coke analysis.^{11,8} The food culture account implies that consumers hold genuine aspirations research must learn to reach through context-sensitive methods. The New Coke account implies that consumers also hold genuine contrary positions, minority signals, that research systematically erases through

aggregation. These framings are theoretically compatible; they imply different methodological remedies. Neither resolves the other, and the field isn't close to synthesis.

Vocabulary That Gained Currency in Q1 2026

Research theater

“outputs that perform the function of research without its epistemological content”

A condition in which research deliverables are produced and consumed as if they constitute evidence, while the underlying process has been degraded to the point where the outputs can't bear the inferential weight placed on them. The term entered Curiosity Current's indexed record in Q4 2025 and was given methodological grounding in Q1 by a synthetic-data platform founder who traced it to the absence of a testable sandbox in consumer research.

Founder, Simulacra Synthetic Data Studio.¹

Selective, not reluctant

“They're not reluctant to use it. They're selective.”

A corrective framing applied to older adults and technology adoption, distinguishing a conversion problem (reluctance requiring persuasion) from a design problem (selectivity requiring better fit). The distinction redefines the research question and the business intervention.

Senior consumer insights manager, AARP.⁴

The guardianship claim

“We need to be guardians of that.”

The emerging professional identity position that researchers carry an analyst’s remit and, beyond it, a duty to defend methodological legitimacy, specifically the distinction between a number and a replicable number, against tool-mediated erosion. The framing appeared with unusual frequency and independence across Q1.

Global head of research and insights, Snap Inc.⁷

Causal scenario modeling

“causal AI that models relationships within existing survey datasets”

The third-generation synthetic data approach that distinguishes itself from LLM-generated respondents by anchoring inference to real human signal. The method extends analysis into underpowered subgroups by modeling conditional relationships within genuine respondent data, rather than generating respondents from training corpora.

Founder, Simulacra Synthetic Data Studio.¹

The just-add-sugar problem

“the regression is gonna tell you to add more sugar... most of the time, these companies can’t add more sugar”

The failure mode in which algorithmic analysis correctly identifies a correlation but can’t apply the domain knowledge that makes acting on that correlation impossible or counterproductive. A category error produced by delegating interpretation to a tool that has no access to formulation constraints, brand positioning, or operational reality.

Founder, Simulacra Synthetic Data Studio.¹

The decision test

“every research question that you are answering should seek to influence a business decision”

A prioritization heuristic that evaluates research projects by their ability to be linked to a specific, identifiable business decision. Questions that fail the test are candidates for deprioritization or elimination.

Group Product Marketing Manager for Insights and Strategy, LinkedIn.¹⁰

The underpowered study

A study designed or commissioned in a manner that renders it statistically incapable of returning actionable results for the subgroups of interest, regardless of execution quality. Q1 surfaced evidence that this represents a substantial fraction of industry output rather than an edge case.

Founder, Simulacra Synthetic Data Studio.¹

Where the Field Is Moving

SCORECARD: Q4 2025 REVISITED

Q4 2025 closed with four falsifiable predictions. One quarter is too short to settle year-horizon calls, but Q1 2026's conversations are the first evidence against which they can be checked. The read below is directional, not conclusive, and where the signal refined a prediction rather than confirming it, that's said plainly.

PARTIALLY HELD

Synthetic respondent practices will face an institutional reckoning, not a market correction.

The epistemological critique intensified and diffused as predicted: Q1's dominant through-line located the "research theater" failure at its methodological root, and the market began separating causal-modeling-on-real-data from LLM-generated respondents into distinct legitimacy tiers.¹ But the specific trigger, a publicly attributable failure shifting the conversation from epistemology to liability, hasn't yet occurred. Q1 forecasts it as still ahead, which leaves the reckoning forming but unrealized.

PARTIALLY HELD

The one-thing discipline will enter practitioner training curricula.

The reduction principle didn't surface in Q1 as a named competency standard. What strengthened instead was the adjacent decision-linkage test, research measured by decisions influenced, not studies completed, converging independently across three organizational contexts with no shared incentive.^{10,7,2} The family of "research must reduce to a decision-relevant claim" is consolidating; the specific codification outcome sits on an end-2026 horizon and isn't yet observable.

NOT YET EVIDENCED

Analog experience research will receive a dedicated investment line in a major consumer-goods category.

Q1's food and consumer signal concerned *what* consumers now value, quality, sourcing, the MAHA/GLP-1 convergence, not a research-budget reallocation toward observational or ethnographic method against digital over-indexing.¹¹ The prediction's falsification horizon is the 2026 planning cycle, which Q1 doesn't reach; this is unconfirmed, not refuted.

PARTIALLY HELD, REFINED

Trust-building will be formalized as a research competency, not a soft skill.

The professionalization signal strengthened materially, but Q1 reframed its axis: what's cohering into an explicit professional stance is *methodological guardianship*, researchers as custodians of replicability, rather than trust-architecture specifically.^{7,5} The underlying thesis held, the profession is codifying what it used to treat as interpersonal luck, but the axis of codification was rigor-defense, not relationship-building. A useful correction to the original framing.

ACCELERATING

- ▲ **Causal AI on real data as the legitimate synthetic tier.** The definitional pressure on “synthetic data” is separating the market into a tier with genuine epistemological legitimacy (causal modeling on real respondent datasets) and one without (LLM-generated respondents). Procurement that doesn't make this distinction is producing false equivalences. The separation will accelerate as quality failures from LLM-based synthetic respondents become attributable.¹
- ▲ **Researcher identity as a defensive professional strategy.** The guardianship framing, researchers as custodians of replicability, has appeared with enough frequency and independence across Q1 to suggest it's cohering into an explicit professional stance. As AI tools continue to reduce the barrier to producing

research-shaped outputs, practitioners with a clear account of what they provide that the tools can't will be better positioned to defend their function.^{7,5}

- ▲ **Quality-dimension tracking in food and consumer categories.** The convergence of MAHA, GLP-1, and progressive food culture consumers around food quality is, if Hartman Group's longitudinal read is accurate, a structural demand shift that volume-focused research designs will systematically miss. Research infrastructure that tracks quality dimensions of purchase decisions will surface signals aggregate volume data can't.¹¹
 - ▲ **The fraud arms race in online sampling.** Bot responses, professional survey farmers, and AI-assisted identity fabrication now represent a systemic threat to the epistemological integrity of online research, well past the scale of a marginal nuisance, and the same capabilities accelerating research synthesis are simultaneously lowering the barrier to fabricating convincing respondents. Sampling quality infrastructure that isn't actively investing in detection and verification is losing ground.⁶
-

SUBSIDING

- ▼ **The volume metric in food and beverage.** As GLP-1 adoption and food quality convergence reshape consumer purchasing, brands and research programs organized around volume as the primary success metric are losing the ability to see the relevant demand signal. Volume tracking misses the quality-driven trade-up that's compensating for reduced intake across converging consumer cohorts.¹¹
- ▼ **Category expertise as a hiring premium.** The counterintuitive hiring argument from corporate insurance, that category-specialist researchers are systematically blind to what's surprising about a category, may have broader applicability. In a fast-changing research environment, the ability to ask fresh questions is worth more than accumulated category familiarity. The premium on insider knowledge is diminishing.²
- ▼ **Qualitative methods as supplementary.** The New Coke failure and the LEGO Serious Play methodology both make versions of the same argument: qualitative signals aren't decoration on quantitative findings, they're the early warning

system. Research architectures that treat qualitative as secondary and aggregation as authoritative are structured to produce the New Coke error repeatedly.^{8,3}

THE YEAR AHEAD

The AI sampling threat will produce a visible quality incident at the industry level, a high-profile research failure attributable to AI-assisted survey fraud, that accelerates the adoption of fraud detection infrastructure and creates procurement pressure for verified panels.

Signal: Multiple Q1 voices, from independent vantage points spanning a panel quality CEO and a food marketing academic, converged on the sampling integrity problem without coordination; the shared concern suggests the industry is approaching a threshold where the gap between fraud prevalence and detection capability becomes publicly visible.^{6,8}

The “synthetic data” label will undergo either market-driven disambiguation or regulatory and industry-body clarification, separating causal scenario modeling from LLM-generated respondents as distinct product categories with distinct quality standards.

Signal: A synthetic-data platform founder spent a significant portion of a Q1 episode correcting his category’s definition instead of advocating for it, a sign that the definitional confusion is causing real market damage severe enough to require active correction.¹

The decision-linkage test for research prioritization will move from practitioner heuristic to institutional standard, with more research functions formally measuring themselves by decisions influenced rather than studies completed.

Signal: Practitioners from LinkedIn, Snap, and New York Life Insurance, three independent organizational contexts with no shared institutional incentive, all converged on decision-influence as the primary measure of research function value in Q1.^{10,7,2}

Q1 2026 Contributors

- [1] Jason Cohen, Founder, Simulacra Synthetic Data Studio. Synthetic data taxonomy, causal AI, LLM respondent critique (Q1 2026). [\[Episode link TK\]](#)

- [2] Charitie Dantis-Gayo, Corporate Vice President, Center for Consumer Insights, New York Life Insurance. Storytelling, hiring philosophy, COVID generation development challenge (Q1 2026). [\[Episode link TK\]](#)

- [3] Garret Westlake, Vice Provost for Innovation and Strategic Design, Virginia Commonwealth University. Play-based research methods, human irreducibility in qualitative inquiry (Q1 2026). [\[Episode link TK\]](#)

- [4] Brittne Kakulla, Senior Consumer Insights Manager, AARP. Older adult technology adoption, universal design, curb-cut principle (Q1 2026). [\[Episode link TK\]](#)

- [5] Nancee Halpin, Founder, Strela. Researcher identity in the AI era, professional differentiation (Q1 2026). [\[Episode link TK\]](#)

- [6] Lisa Wilding Brown, CEO, Innovate MR. Panel quality, AI-accelerated survey fraud, sampling integrity (Q1 2026). [\[Episode link TK\]](#)

- [7] Aarti Bhaskaran, Global Head of Research and Insights, Snap Inc. Methodological guardianship, AI in research workflows, professional voice (Q1 2026). [\[Episode link TK\]](#)

- [8] Ernest Baskin, Department Chair and Associate Professor of Food Marketing, Saint Joseph's University. Pre-registration, qualitative-quantitative integration, New Coke (Q1 2026). [\[Episode link TK\]](#)

- [9] Christopher Khoury, Vice President of Strategic Insights, American Medical Association. Healthcare stakeholder complexity, organizational navigation (Q1 2026). [\[Episode link TK\]](#)

- [10] Kylee Lessard, Group Product Marketing Manager for Insights and Strategy, LinkedIn. Research-business decision linkage, strategic function positioning (Q1 2026). [\[Episode link TK\]](#)

- [11] Laurie Demeritt, CEO, The Hartman Group. Food culture, say-do gap reframe, MAHA/GLP-1/progressive convergence, trend identification (Q1 2026). *[Episode link TK]*

About This Paper

Indexed record as instrument. This white paper treats the indexed record of eleven Q1 2026 episodes of the Curiosity Current podcast (the *corpus*) as the instrument of analysis. The unit of analysis is the phenomenon, the recurring pattern, emergent tension, or convergent observation, not the individual guest. Practitioners appear as data points; their arguments are examined for convergence, divergence, and explanatory weight against the quarter's through-lines.

Attribution format. In-body attribution uses role and organization where the source's institutional vantage point carries analytical weight. Full personal names appear only in References & Contributors. Superscript numerals are assigned in order of first appearance and reused on subsequent citations of the same source.

Verification flags. Two statistical claims are flagged for independent verification before use as cited evidence:

- The assertion that more than 50% of surveys return no statistically valid results (attributed to the founder of Simulacra Synthetic Data Studio).¹
- The claim that approximately 1 in 2 clicks online today isn't generated by a real human (attributed to the CEO of Innovate MR).⁶

Both claims are attributed clearly to their sources. Neither should be treated as independently verified fact. Their directional force is noted where relevant; specific numeric precision should be independently confirmed.

Cross-references. Convergences and tensions noted in this paper are drawn from the Curiosity Current indexed episode record. Q1 2026 introduces eleven contributors. The "research theater" cross-reference to Q4 2025 is drawn from the Q4 paper's indexed episodes.

Quarterly scope. All references to episode counts, contributor counts, and cross-reference counts in this paper are scoped to Q1 2026 and don't represent living totals from the full Curiosity Current indexed record.

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

Curiosity Current · AYTМ · Q1 2026

Compiled by the Curiosity Current editorial team · Published Q2 2026



© 2026 aytм. All rights reserved.