

Not Deception, But Design: The Umbra Cone And Structural Risk-Shadowing in AI-Mediated Commerce Extending a Hybrid-Domain Theory of Meaning-Making from Science to Commerce

Luca Magni*

Professor of Practice Luiss Business School Villa Blanc, Via Nomentana 216, 00162 Rome, Italy

***Corresponding Author:** Luca Magni, Professor of Practice Luiss Business School Villa Blanc, Via Nomentana 216, 00162 Rome, Italy.

Submitted: 23 July 2026

Accepted: 28 July 2026

Published: 06 Aug 2026

Citation: Luca Magni (2026). Not Deception, But Design: The Umbra Cone And Structural Risk-Shadowing in AI-Mediated Commerce Extending a Hybrid-Domain Theory of Meaning-Making from Science to Commerce. *American Journal of Marketing, Finance and Trade*, 1(1), 01-09.

Abstract

In the winter of 2023, Coca-Cola released a holiday advertisement remade with generative artificial intelligence, confident in the technology's efficiency and polish [1]. Viewers called it "soulless," and the backlash became a small industry parable: fluency is not the same thing as trust. This article argues that the parable generalizes, and grounds that argument in existing empirical and legal evidence rather than theory alone. Marketing copy, financial advisory scripts, and cross-border trade correspondence are all now drafted in part by generative AI, and several independent research literatures document the same underlying pattern with real data: reward-based crowdfunding campaigns that lean on an overly warm AI disclosure see funding penalties worsen through measured increases in AI washing suspicion; earnings-call analyses show confident AI talk moving stock prices well before any AI walk, measured through patents and hiring, catches up; and finance research documents markets rewarding unsubstantiated AI narratives before eventually, and measurably, punishing them. What none of these literatures offers is a shared mechanism for why the gap between talk and walk keeps appearing. This article proposes one, extending the Umbra Cone, a Learnable Theory construct built to explain a strikingly similar pattern in AI-mediated science communication, from science to commerce. The claim is that AI-mediated trade communication systematically foregrounds belonging and aspiration while quietly dimming the salience of risk, limitation, and fine print, a structural byproduct of how generative systems optimize for fluency rather than a matter of managerial intent, and one that a well-documented psychological tendency, the machine heuristic, keeps readers from catching on their own. Five testable propositions follow, most operationalized with already-validated tools rather than the theory's own pilot instrument alone, and the regulatory discussion is anchored in the United States Supreme Court's 2024 ruling distinguishing actionable half-truths from unactionable pure omissions under securities law.

Keywords: Learnable Theory, Umbra Cone, Ai-Mediated Communication, AI washing, trade communication, marketing communication, Complementary Semantic Distance Index

Introduction

Picture two scenes, a year and a world apart. In the first, a marketing team at a major beverage company unveils a holiday commercial remade almost entirely by generative AI, proud of what the technology let them do faster and cheaper than ever before. Audiences recoil, and one word recurs across the coverage: viewers dismissed the ad as "soulless" despite the company's own enthusiasm for what the technology had achieved [1]. In the second scene, a factory manager in Shenzhen is closing a sample order with a buyer in Mexico, the two of

them separated by language and twelve time zones, brought together by an AI-powered trade platform that translates not just his words but his intent, adapting technical phrasing to local convention in real time. As the general manager overseeing the deal put it, "the efficiency is unimaginable compared to relying on traditional translation methods" [2]. Both scenes involve the same technology solving, on the surface, the same kind of problem: making commercial communication faster, cheaper, and more fluent. Only one of them backfires. The question this article asks is why, and whether the answer is something more

principled than an anecdote about tone.

The backlash and the rapport are not opposites so much as two faces of the same coin. By 2024, every marketing professional surveyed in one industry study reported using generative AI in some part of their work, seduced by its efficiency and the apparent quality of its output [1]. Yet the same years produced a stubbornly widening gap between how advertising executives believe their AI-generated campaigns land and how audiences actually feel about them, a gap that grew rather than closed even as the technology matured [3]. Financial markets tell a version of the same story with real money attached: firms that lean hard into AI narrative in their disclosures get an early reward from investors, a flicker of enthusiasm, and then, once the gap between what was claimed and what was actually built becomes visible, a harder correction than if they had said nothing grand at all [4]. Regulators have already given this pattern a name, AI washing, described as firms creating a “façade of innovation and reliability” for investors and customers [5]. and already delivered a first scalp: in March 2024 the United States Securities and Exchange Commission (SEC) charged two investment advisers with dressing up ordinary software in AI language their products did not earn [5].

What is curious is that none of these episodes involves an outright lie in the crude sense. The AI disclosure is usually there, somewhere, in the terms of service. The risk warning is filed, technically, in the prospectus. The model’s limitations are documented, technically, in an appendix nobody reads before signing. Nothing essential is missing from the record. What changes, case after case, is not the presence of the caveat but its claim on the reader’s attention, crowded out by a register built to sound confident, warm, and fluent. This article argues that this is not three unrelated stories about disclosure, trust, and fraud. It is one story, told from three vantage points, and Learnable Theory already has a name for the mechanism doing the crowding out: the Umbra Cone.

The Umbra Cone was first built to explain a nearly identical asymmetry in a different setting, AI-generated explanations of scientific findings, where it was shown that a confident, fluent explanatory voice tends to shadow the methodological caveats a claim’s real epistemic weight depends on [11]. Science communication was chosen as the theory’s first home because the shadowed content there is unusually easy to pin down: a sample size either was reported or it was not, a confidence interval either survives translation into plain language or it does not. Commerce is messier, but the stakes of what goes missing in the shadows are arguably higher, a fiduciary caveat, a liability clause, an undisclosed error rate carry consequences that are financial and legal, not merely epistemic. If the Umbra Cone names something real about how AI-generated fluency behaves, rather than something specific to how scientists happen to write, it ought to show up wherever that fluency is deployed. This article follows that thread into trade.

What follows moves in five steps, and it leans throughout on existing empirical and legal work rather than asking the theory

to stand on its own. Section 2 walks through the Learnable Theory apparatus needed to make the extension legible: the Learnable and learnables, the LEBO strata, the Conic Helicoid, the Umbra Cone itself, the Complementary Semantic Distance Index, and Differenza Inversa. Section 3 gathers the marketing, fintech, and cross-border trade literatures that, read separately, have each brushed up against this problem without naming it, including recent work quantifying the gap between what firms say about their AI and what their hiring and patent records actually show. Section 4 makes the extension explicit, arguing for a need-domain asymmetry at the heart of AI-mediated commercial fluency, distinguishing that asymmetry from older accounts of deliberate managerial deception, and identifying when it should recede rather than deepen. Section 5 turns the argument into five falsifiable propositions, most of which can be tested with existing, already-validated tools rather than the theory’s own pilot instrument alone. Section 6 asks what any of this is good for, grounding the regulatory discussion in a 2024 Supreme Court decision on securities half-truths rather than treating disclosure policy in the abstract. Section 7 is honest about what this article has not done.

Theoretical Background: Learnable Theory and the Hybrid-Domain Account

Learnable Theory grew out of three founding texts and has since been extended into neurosemantics, organizational alignment, and AI-mediated science communication. Rather than rehearse the whole architecture, this section introduces only the pieces the rest of the argument leans on, in roughly the order they will be needed [6,7,8,9,10].

The Learnable and Learnables

Start with a distinction that sounds almost too simple to matter and turns out to matter quite a lot. The Learnable, capitalized, is the universal cognitive-affective horizon within which any meaning-making happens at all, the shared ground beneath every particular way of making sense of things. A learnable, lowercase, plural, is one of those particular ways: a brand’s house voice, a regulator’s disclosure template, a negotiation script handed down through a trade desk. A learnable is one constituent framework within the Learnable, never a synonym for it [5]. Every advertisement, every advisory email, every negotiation message is the deployment of a specific learnable against that wider backdrop, and it is worth holding onto this distinction, because much of what follows is really a story about what a given learnable chooses to make salient and what it lets slip into the background.

Lebo Stratification

Borrowing from Bhaskarian critical realism and extending it with a neurosemantic layer, LEBO separates four strata of communicative reality: the Real, what is true independent of anyone saying so, an actual error rate, an actual contractual obligation; the Actual, what specifically gets said or written; the Empirical, what a reader or listener actually registers; and the Learnable, the interpretive frame through which that registering happens [10]. The whole argument of this article lives in the gap between the Actual and the Empirical. The shadowed content,

in every case examined below, is fully present at the Actual level, sitting right there in the text, and yet fails to register at the Empirical level with anything like the weight its Real-world importance would justify. Nothing is deleted. Something simply does not land.

The Conic Helicoid and the question of who is answerable

Picture meaning-making not as a straight line but as an ascent around a cone, each turn revisiting what came before while adding a new layer of sense, anchored by an apex, a committed point of view someone is willing to be held to. Take away the apex and the ascent becomes degenerate, circling the same material without ever resolving into a stance anyone stakes their name on. This matters later, in Section 4.5, because an apex-anchored review, a named professional whose license or reputation is on the line, turns out to be the theory's best candidate for what stops shadowing from happening in the first place.

The Umbra Cone: The Shadow Every Claim Casts

Every assertion, hedged or not, casts a shadow, a region of unstated uncertainty or unresolved boundary condition that surrounds it. A hedge is really just an admission that part of that shadow exists. What turns out to matter most, though, is not how big the shadow is but where its apex sits, whether the uncertainty is anchored to someone specific who owns it, or left ambient and ownerless, floating free of any accountable voice. In its extension to AI-mediated science communication, this idea sharpened into what was called the umbra cone effect: certain domains of concern get disproportionately shadowed by a given register, leaving readers with a false sense of certainty even when the caveats were, technically, right there on the page. In that setting, the domain that kept vanishing into shadow was methodological limitation, a Safety-domain concern in the theory's own vocabulary, eclipsed by prose that sounded too confident to invite doubt [11]. Section 4 argues that trade communication runs the identical trick, only the specific thing hiding in the shadow changes with the setting.

The Complementary Semantic Distance Index

If the Umbra Cone names the phenomenon, the Complementary Semantic Distance Index, or CSDI, is the attempt to measure it. CSDI sorts communicative content into four need-domains, Survival, Safety, Sociality, and Self-Actualization, each carrying a Positive, Neutral, or Negative valence, and computes how far apart two passages sit in that space [7]. A register whose foregrounded content lands in one corner of that space while its shadowed content sits in a distant corner shows a large CSDI distance; a register where the two stay close shows a small one. It remains, honestly, a pilot-plausibility tool rather than a fully validated instrument, and that honesty carries through explicitly into the limitations discussed later.

Differenza Inversa And Prepositional Focusing

Differenza Inversa, inverse difference, describes something quieter still: a compensating movement in which a visible gain along one dimension is offset by an invisible loss along another, first documented as an intrapsychic operator by which a

remembered event's emotional charge can flip without the facts themselves changing. Prepositional focusing and defocusing, the foregrounding of one element and the relegation of specific others performed by prepositions and other relators, are one instance of the broader family of attentional mechanics that make such shifts possible in language, distinct from the focusing effects proper to other lexical classes [12].

Why A Theory Built for Science Should Also Explain Commerce

The hybrid-domain account underlying all of this treats AI-authored text as sitting in a third space, neither the genuine, agent-like authorship a strong emergentist view would claim, nor the neutral, contentless mirror an instrumentalist view would expect [11]. Neither extreme actually predicts what gets observed. If the AI were truly an author with something like intent, its shadowing should look arbitrary, idiosyncratic, tied to whatever the model happens to care about. If it were truly just a mirror, there should be no systematic shadowing pattern at all, only noise reflecting whatever bias sits buried in its training data. What is observed instead looks patterned, and the hybrid-domain account locates that pattern in the generative architecture itself, in how these systems allocate fluency and confidence across a response, largely regardless of what the response happens to be about. If that diagnosis holds, the same architecture drafting a scientific summary, a marketing email, or a trade contract should shadow in the same structural way, even as the specific thing being shadowed changes with the room it is standing in.

The Reader's Half of The Mechanism: The Machine Heuristic

A theory of shadowing built entirely around what the generative architecture does would still be incomplete, because a shadow only works if the reader does not compensate for it, and there is no obvious reason a reader should fail to compensate unless something about the source of the text lowers their guard in the first place. This is where a second, independent literature does foundational rather than merely supporting work. Research on the machine heuristic has shown for well over a decade that people extend a specific kind of trust to output they believe came from a machine rather than a person, assuming it to be more objective, more accurate, and less self-interested than the equivalent human output, an assumption that itself does much of the persuading before any actual evaluation of content takes place ([13,14]. A human salesperson who sounds a little too confident tends to raise a reader's guard. A machine that sounds equally confident tends to lower it, precisely because the reader has already extended the machine a presumption of objectivity a human speaker would have had to earn first.

Built into the theory from this point forward, then, is a two-sided account rather than a one-sided one. The Umbra Cone names what the generative architecture does: foreground fluency, let qualification recede. The machine heuristic names what the reader does in response, or rather fails to do: apply the skepticism they would ordinarily reserve for an equally confident human voice. Section 4 develops the commercial version of the architecture's

half of this mechanism in detail; the reader's half, established here, is what makes that architecture's output land as intended rather than triggering the doubt a human speaker using the same words would invite.

Three Rooms, One Draft: A Literature Review

The Marketing Floor

Generative AI has moved, in the space of a couple of years, from a novelty to the default way a great deal of marketing copy gets written, used to draft ads, personalize messages, and speed up creative work that used to take weeks [1]. The Coca-Cola episode that opened this article was not a fluke. Systematic reviews of AI-generated advertising find that disclosing AI authorship measurably reduces trust and softens attitudes toward an ad, and the effect bites hardest precisely where a brand is trying to say something intangible about itself rather than list a product spec [15]. Head-to-head comparisons of generative AI and human-written brand narratives find the gap is not about logic, AI holds its own there, but about feeling, about the emotionally attuned register that still seems to elude machine-generated prose [16]. Whether and how to disclose AI's hand in a piece of content has become, almost overnight, a design decision with real consequences for trust [17]. And the industry's own tracking suggests nobody inside the room has quite noticed how the audience actually feels: the gap between what advertising executives believe consumers think of AI-generated ads and what consumers actually report has widened, not narrowed, growing from thirty-two points to thirty-seven between 2024 and 2026 [3].

A more fine-grained empirical picture comes from reward-based crowdfunding, where individual disclosures and individual funding outcomes can be matched one to one. Using Kickstarter's mandatory AI disclosure policy as a natural experiment, recent research finds that disclosed AI involvement cuts funds raised by nearly forty percent and backer counts by close to a quarter, but the size of that penalty depends heavily on how the disclosure is written [18]. Campaigns that lean on higher explicitness and authenticity see the penalty soften. Campaigns that lean instead on an excessively positive emotional tone, the instinctive move for a creator trying to reassure backers that everything is fine, see the penalty worsen, and the same study traces exactly why: an overly warm register raises AI washing suspicion and lowers perceived creator competence, and it is those two mechanisms, not the AI disclosure itself, that account for the drop in pledges [18]. That is a strikingly literal, already-measured instance of a register foregrounding warmth at the direct expense of the very credibility it was trying to protect.

The Advisory Desk

Financial advice has been quietly automated for long enough that robo-advisors, platforms that gather a client's risk profile and hand back a recommended portfolio without a human ever entering the conversation, are no longer the experiment they once were [19]. That automation lands on a client base that was never uniform to begin with; research on traditional human advice already found that no single advisory model fits every

investor equally well [20], which raises the stakes of whatever tone an automated substitute adopts in its place. The trouble is not merely that the algorithms are hard to explain to a regulator, though they are, and increasingly that difficulty is treated as a compliance problem in its own right rather than a temporary technical hiccup. The deeper trouble is what the finance literature has started calling AI washing outright, the systematic gap between what a firm's disclosures claim about its AI and what its actual deployment can support [21]. Using patent activity as a rough proxy for real capability and text analysis of the disclosures themselves, this research finds a now-familiar shape: unsubstantiated AI narratives "initially attract investors but ultimately yield negative market reactions"), rewarding the confident story first and punishing it only once the underlying substance becomes visible. The SEC's own 2024 enforcement action against two advisers making AI claims their software could not back up reads like a case study lifted straight from that findin [5,21].

Two further literatures, read together, give this backlash both a name and a measurement strategy. On the talk side, large language models applied to earnings-call transcripts can now quantify exactly how much a firm discusses generative AI and how that discussion trends over time; this research finds that discussion itself moves markets, with a one percentage point increase in GenAI exposure associated with a measurable rise in quarterly excess stock returns, well before any of that discussion has translated into a documented change in expected earnings [22]. On the walk side, a separate and considerably larger literature has learned to measure the substance behind the talk directly, using firms' own hiring records and patent filings rather than what executives choose to say about themselves; this walk-side measure finds that firms genuinely investing in AI do experience real gains in growth, product innovation, and valuation [23]. Talk and walk, in other words, are each independently measurable and each independently real. What the AI washing literature adds is the observation that they do not have to move together, and that markets eventually notice, and punish, the firms for whom they do not [4].

The Negotiation Table, Twelve Time Zones Wide

The third room is cross-border trade, where AI increasingly sits between a buyer and a seller who do not share a language at all. Industry accounts describe platforms that do more than translate vocabulary, adapting phrasing to local convention in the moment, with operators reporting that this shortens the distance between a first inquiry and an actual purchase intent, echoing the Shenzhen manager's story that opened this article [2]. It has to be said plainly that this specific claim, that generative AI mediation measurably increases fluency and rapport in live cross-border trade negotiation, comes almost entirely from trade press and platform marketing rather than peer-reviewed research, and no amount of adjacent literature closes that particular gap. What adjacent literature does exist, and it is worth taking seriously rather than ignoring, concerns not live negotiation but the written legal and contractual documents such negotiations eventually produce. Studies assessing machine translation against legal

and contractual terminology find accuracy dropping sharply once context sensitivity is required, from roughly two-thirds accuracy on straightforward terms to just over half once a term's meaning depends on its surrounding legal context [24]. A direct comparison of human and AI translation of legal documents, including commercial contracts and agreements, found human translators significantly outperforming AI across terminology accuracy, clarity, and adherence to the relevant legal framework [25]. Neither study is about generative AI mediating a live negotiation, and neither should be read as if it were; both are about translation quality on static legal text, a related but distinct task. What they establish, credibly and with peer review behind them, is that automated translation has a documented, measurable tendency to lose exactly the kind of contextual and technical precision a contract depends on, which is the premise Proposition 4 needs and the trade press evidence alone could not supply.

Three Rooms, One Draft

Put side by side, these three literatures are describing the same shape from three different windows. Marketing research shows AI-generated content underperforming on trust and emotional register relative to how confident it sounds. Financial disclosure research shows AI narrative and AI substance drifting apart in ways markets eventually, if not immediately, notice and punish. Trade communication accounts describe fluency gains without anyone systematically checking whether technical precision kept pace. None of these three rooms has a name for what they are all, separately, circling. That absence is the gap this article tries to close.

The Umbra Cone Walks into a Trade Fair A Baseline Foil: This is Not Stein's Manager

Before stating the central claim, it is worth being explicit about what this article is not arguing, because the alternative explanation is well established and deserves to be named up front rather than fended off later. Economists have had a working model of confident corporate overclaiming for decades. Stein's classical account of myopic corporate behavior describes a manager who forsakes genuinely good investments in favor of choices that inflate the market's short-term read on the firm's worth, fully aware of what he is doing and why [26]. That manager has intent: he knows the earnings figure is inflated and chooses to inflate it anyway, gambling that the market will believe him long enough for it to matter. If the Umbra Cone in commerce were nothing more than Stein's myopic manager wearing a chatbot as a mask, the construct would not be adding much to a story economics has already told well.

It is not that story, and the difference is the foundation everything else in this section builds on. A generative model trained to maximize fluency, engagement, and reader satisfaction will tend to foreground warm, confident, Sociality-rich language and let drier Safety-domain qualification recede, not because anyone instructed it to conceal anything, but because that is what a system optimized for those objectives does with language by default. An executive with entirely good intentions, asking an

enterprise assistant to draft a balanced risk disclosure, can still receive back a draft with the Umbra Cone effect already built into it, simply because confident phrasing is what the underlying model has learned readers respond to. The shadowing this article describes is architectural rather than intentional, which is exactly why disclosure rules built to catch intentional deception, rules that ask whether a specific statement is true or false, are not well suited to catching it. And it survives on the reader's side of the exchange for the reason established in Section 2.8: the machine heuristic lowers the very guard that would ordinarily catch an equally confident human voice, which is what lets an unintentional, architectural shadow do the persuasive work that would otherwise require Stein's manager and his deliberate choice to lie.

The Need-Domain Asymmetry Hypothesis

With that foil in place, here is the central claim, stated plainly: AI-mediated trade communication, whether it is selling a product, managing a portfolio, or closing a cross-border deal, will tend to foreground belonging, aspiration, and confidence, the Sociality and Self-Actualization domains in CSDI terms, while letting risk, limitation, and fine print, the Safety domain, fade into the background. Nothing about this requires the AI to omit anything. The mechanism is need-domain foregrounding: the register stays technically complete, every caveat present somewhere in the text, but the fluency, the confidence markers, and the narrative momentum all get spent on the domains that sell, leaving the domains that protect the reader sitting quietly, unread, in the background.

Three Rooms, Three Shadows

In marketing, what usually falls into shadow is the fact of AI authorship itself, along with the data practices quietly underwriting all that personalization. The moment that shadowed content gets deliberately dragged into the light through disclosure, trust drops, which is itself indirect proof that the shadow was doing real persuasive work while it lasted [15]. The Kickstarter evidence reviewed above shows the same mechanism at the level of a single disclosure: a register dialed toward Sociality-Positive warmth, meant to reassure a backer that a human stands behind the work, instead draws attention to the very Safety-domain question, whether that warmth is compensating for a lack of genuine human involvement, that it was trying to keep out of view, and the resulting drop in perceived competence and rise in AI washing suspicion are measured outcomes, not a theoretical inference [27]. Financial advisory is a more tangled case, because the Safety domain splits in two: a Safety-Positive register, all security and confidence and protection, gets pushed to the front, while a Safety-Neutral or Safety-Negative register, the actual model limitations, the actual error rates, gets pushed to the back. That split is exactly what the talk and walk literatures are measuring from two different angles, a short-term reward for the confident talk, a longer-term penalty once the walk-side substance, or its absence, becomes visible [22]. At the negotiation table, the proposed casualty is technical and contractual precision, the specification limit, the liability clause, the delivery contingency, smoothed away by a

translation layer optimized for rapport rather than for preserving exactly where a claim's boundary sits.

A Borrowed Idea, Worn Carefully: Differenza Inversa at the Negotiation Table

There is a further step worth taking, offered with the same caution the theory applies whenever a construct is asked to do work outside its home turf. Differenza Inversa describes a compensating movement, a visible gain in one place quietly offset by an invisible loss somewhere else. Applied here, purely as an analogy, the proposal is that the fluency and rapport a buyer or seller consciously notices and credits to an AI-mediated negotiation, the visible Sociality gain, may travel alongside an invisible dimming of technical and contractual specificity, a Safety-domain cost that nobody in the room consciously registers while the deal is being struck. The two situations are not the same mechanism wearing different clothes: the original construct describes a reversal of viewpoint, this one describes a reallocation of salience. But the compensating shape, gain here, quiet loss there, is recognizable enough to be worth testing directly, which is what Proposition 4 sets out to do. It is also, of the four commercial extensions in this article, the one with the thinnest direct empirical floor, a gap Section 3.3 already flagged and Proposition 4 below addresses head on rather than papering over.

When the Shadow Lifts

The Umbra Cone was never meant to describe an inescapable fate. Its own logic ties the severity of shadowing to whether the shadow's apex is anchored to someone accountable or left to drift. Translated into commerce, this predicts that shadowing should soften wherever AI-generated commercial text passes through a documented human check tied to a specific person's professional accountability, an apex-anchored step in Conic Helicoid terms, rather than going straight from model output to client inbox. Some accounts of AI deployment in regulated finance already describe pipelines built exactly this way, with a named, accountable reviewer standing between the model and the client, a pattern regulators appear to be treating as a baseline expectation rather than a nice-to-have [31]. The theory's prediction, formalized as Proposition 5 below, is that such pipelines should show a measurably smaller gap between what gets foregrounded and what gets shadowed than pipelines without that check, no matter how much of the drafting the AI actually did.

Five Propositions, Offered as an Invitation

None of what follows should be read as a finding. These are hypotheses stated precisely enough to be tested, in the same spirit as the propositions offered when this theory was first extended to persuasive hedging [12].

P1 (Marketing shadowing). AI-generated marketing content that does not disclose its AI authorship will show a larger CSDI distance between its foregrounded need-domain configuration, expected to be Sociality-Positive or Self-Actualization-Positive, and its shadowed configuration, expected to be Safety-Neutral,

than human-authored content matched for product category and campaign objective. A weaker, immediately testable version of this proposition already has indirect support: disclosed campaigns that lean into an excessively positive emotional tone see funding outcomes worsen specifically through heightened AI washing suspicion and reduced perceived creator competence, consistent with Sociality overforegrounding drawing attention to, rather than away from, the shadowed Safety-domain question of authorship [22,27].

P2 (Disclosure backlash as a correction). Firms whose AI disclosures show a larger gap between talk, AI narrative intensity measured with the sentometric approach already validated for earnings-call GenAI exposure, and walk, independently verifiable AI capability measured through resume-based or patent-based indices, will suffer larger negative abnormal returns once that gap becomes visible to the market, than firms whose talk and walk already sat close together. Because both the talk-side and walk-side measures are already established, published, and openly documented tools, this proposition is directly testable without reference to CSDI [23].

P3 (Advisory shadowing). Robo-advisory communications thick with Safety-Positive language, confidence, security, protection, will show correspondingly heavier shadowing of Safety-Neutral content such as explainability limits and error rates, and this gap will shrink measurably under stricter, independently audited disclosure regimes. This is testable through a standard behavioral paradigm adapted from the truth-falsity crossover effect documented in AI-disclosure research[32]: if labeling a financial script as AI-generated inflates the perceived credibility of a version with heavily shadowed caveats while depressing the credibility of an otherwise equivalent version with prominent caveats, that crossover pattern is itself behavioral evidence of the Umbra Cone operating in a financial advisory setting, detectable with a conventional experimental design rather than the CSDI instrument.

P4 (Negotiation shadowing). AI-mediated cross-border negotiations showing a larger fluency and rapport gain over a human-mediated baseline, matched for deal complexity, will show a correspondingly measurable drop in explicit technical or liability language, consistent with a Differenza Inversa compensating pattern applied here as analogy rather than as a documented instance of the source construct. This is the one proposition in this article for which no direct empirical anchor yet exists, and it is presented as such rather than dressed up to look otherwise. The adjacent literature on machine translation of legal and contractual text supports the second half of the claim, that automated language processing measurably loses contextual and technical precision relative to human output, but that literature concerns static document translation, not live generative negotiation, and does not itself establish the fluency-and-rapport gain the proposition's first half predicts. Here is where the empirical data currently stops: no published study yet measures both halves of this trade-off in the same generative, AI-mediated negotiation setting. What the theory

predicts, and what this proposition exists to make testable once that data becomes available, is that the two halves will move together, a rise in sentiment-measured rapport alongside a fall in liability-clause density in the same transcripts. This can be operationalized without CSDI: standard sentiment analysis can capture the fluency and rapport gain as a rise in positive, collaborative sentiment, while automated liability-clause and specification extraction, already routine tools in contract analytics, can independently capture any corresponding drop in explicit technical or liability language across the same transcripts. A measurable rise in the first alongside a measurable fall in the second is the signature this proposition predicts, and neither half of that signature depends on the pilot-plausibility CSDI instrument [1,24].

P5 (Apex-anchoring lifts the shadow). Commercial AI communication passed through a documented, individually accountable human review step before release will show a smaller CSDI distance between foreground and shadow than equivalent communication released without that step, regardless of how much of the drafting the AI actually did.

What This is Good For For Regulators

Most current rules treat AI disclosure as a light switch, on or off: did anyone tell the reader AI was involved. The Umbra Cone suggests the switch is the wrong instrument. A disclosure can be technically on and still be shadowed into irrelevance by everything around it, satisfying the letter of a labeling rule while doing almost nothing to correct the imbalance regulators actually care about. If Proposition 2 holds up, a CSDI-style distance measure would give supervisors something closer to a dimmer than a switch, a graduated read on how far a firm's narrative has drifted from its substance, computable in principle from the disclosure text itself against outside indicators like patent activity, without anyone needing to see inside the model.

Current securities law already draws a distinction this dimmer-switch metaphor can be anchored to. In 2024 the United States Supreme Court held, unanimously, that Rule 10b-5(b) does not reach pure omissions, information simply left out, but does reach half-truths, statements rendered misleading by what surrounds them [26]. Shadowed content, as this article has defined it, is never a pure omission; it sits in the text, satisfying a narrow reading of a disclosure requirement, while being rendered functionally invisible by the register around it. Under the Court's own framing, that is a half-truth, not a gap, and half-truths are the one thing Rule 10b-5(b) was already built to reach. Read this way, the Umbra Cone is less a call for new law than a proposal for operationalizing a distinction the Supreme Court has already drawn, giving the idea of a half-truth a number rather than leaving it to a court's after-the-fact impression of a document's overall tone.

The SEC's 2022 action against BNY Mellon Investment Adviser, though it concerns ESG claims rather than AI, previews how this could work. the firm was not penalized for an outright

fabrication but for representing that a specific review process had been applied to all its holdings when it had not always been applied, a procedural half-truth the Commission treated exactly as seriously as a factual one [30]. The same logic applied to AI narrative would ask not whether a firm's AI claim was technically false, but whether its overall disclosure, narrative and substance together, left a reasonable investor with a materially misleading impression, precisely the question a CSDI-style distance score is built to answer.

The most self-contained remedy this article can offer follows directly from Proposition 5 rather than from outside doctrine. If apex-anchored human review measurably narrows the gap between what gets foregrounded and what gets shadowed, as that proposition predicts, the more immediate regulatory lever is not a new legal theory at all but a compliance requirement already implicit in existing disclosure law: named, individually accountable review of AI-drafted client communication before release, the same apex-anchoring the Umbra Cone construct already treats as the condition under which shadowing recedes. That recommendation stays entirely within the paper's own theoretical borders and does not require importing a remedy built for a different area of law.

A further question, more speculative and offered strictly as a conceptual analogy rather than a legal argument, is whether disclosure regulators might eventually need a remedy that reaches the generative process itself rather than only the disclosure it produced. Outside securities law entirely, the Federal Trade Commission (FTC) has established a remedy sometimes called algorithmic disgorgement, ordering a firm to delete not just unlawfully obtained data but any model or algorithm trained on it, applied in its settlements with Everalbum, Weight Watchers and Kurbo, and Rite Aid [29,30,31]. That remedy sits in a different legal universe from the SEC's half-truth doctrine: FTC disgorgement responds to a model built on tainted data, a data-provenance problem, while the shadowing this article describes is a structural, architecture-level tendency present even when every input was lawfully obtained, a disclosure-integrity problem. The two paradigms are not interchangeable, and this article does not argue that a securities regulator could or should order a firm to delete a language model the way the FTC orders deletion of a facial-recognition model built on scraped photos. What the FTC's willingness to reach the model itself, rather than stopping at the output, does offer is a conceptual blueprint: evidence that a regulator can treat a generative artifact, not just a single disclosure, as the object in need of remedy when the two are structurally entangled. Whether securities regulators will eventually need an equivalent tool for generative language systems is an open question this article raises rather than answers, and it is raised here explicitly as future work, not as a claim about what current law already permits.

For Marketers and Traders

For anyone actually writing this material, the implication is not to strip AI narrative out of the copy but to close the distance between the story and the thing the story is about. The marketing

research already shows that disclosure alone can dent trust when the underlying gap is real; it is a short step, though not yet a tested one, to expect that disclosure paired with a genuinely smaller gap should not cost the same. For platforms running AI-mediated cross-border negotiation, Proposition 4 points to a concrete audit habit worth adopting now rather than later: any system tuned to maximize fluency and rapport should be checked, specifically and separately, for whether it is quietly costing the deal some technical or contractual precision, rather than assuming a gain on one side means no cost on the other.

For the Theory itself

The real payoff here, for Learnable Theory's own research programme, is not a new construct but evidence that an old one travels. The Umbra Cone was built and tested in a single room, science communication, where the thing hiding in the shadow is unusually easy to name. Finding it recognizable in commerce, even at the purely theoretical level attempted here, suggests it is describing something about AI-generated fluency itself rather than something peculiar to how scientists talk, which is exactly what the hybrid-domain account predicted all along. That opens the door to the same question in rooms this article has not entered, legal AI and health communication being the two most obvious next stops, each needing its own translation of what counts as the thing worth protecting from the shadows.

What This Article Has Not Done

This is a work of theory, not of data, and its five propositions are an invitation rather than a result. A few honest caveats deserve to sit here rather than be smoothed away by the same fluency this article is warning against. The CSDI instrument behind Proposition 1 and Proposition 5 is still, as of this writing, a pilot-plausibility tool, so any claim about need-domain distances in commercial text is a hypothesis awaiting its own validation in this new setting, not a settled measurement. The talk-side and walk-side measures behind Proposition 2 and the crossover-effect paradigm behind Proposition 3 are independently validated, published instruments, but none of them has yet been applied to the specific claim this article makes about them; borrowing a validated tool for a new hypothesis is not the same as having already tested that hypothesis with it. The evidence on cross-border AI-mediated negotiation in Section 3.3 comes mostly from trade press rather than peer-reviewed research, a real gap Proposition 4 is designed to help close rather than one this article can claim to have already closed. The extension of *Differenza Inversa* to negotiation fluency in Section 4.4 is an analogy, offered with the same caution applied when the construct first crossed into persuasive hedging, and should not be mistaken for a demonstrated instance of the original mechanism. The discussion of algorithmic disgorgement in Section 6.1 is explicitly flagged there as an open question rather than a legal argument this article has made; the remedy has real precedent, but its extension to AI washing has not been tested in any court or agency proceeding, and this article does not claim otherwise. And finally, there is no new data anywhere in these pages. Whatever this article is worth, it should be judged on whether its propositions are testable and worth testing, not on

anything it claims to have already found [7,22,23,18,12].

Closing Remarks

Go back to the two scenes this article opened with. One is a boardroom watching an expensive ad fall flat, the other is a factory floor in Shenzhen closing a deal in minutes with a buyer half a world away. Both ran on the same underlying technology. The difference between them, this article has argued, is not luck but shadow, whether the fluency doing the persuading left the reader's Safety-domain concerns intact and visible, or quietly dimmed them in favor of a better story. The Umbra Cone, built to explain why AI-generated science writing loses its own caveats in translation, extends without much strain to explain why AI-generated trade writing loses its own fine print the same way. Whether the five propositions above survive contact with real data is, honestly, the more interesting question than anything settled in this article. But the theoretical wager underneath all of it is a simple one: fluency is manufactured by the same machinery no matter what it is being fluent about, and if the Umbra Cone is a property of that machinery rather than a quirk of how scientists write, then commerce was always going to inherit it too.

References

1. NIM (Nuremberg Institute for Market Decisions) (2026) Transparency without trust: Consumer attitudes toward AI-generated marketing content [Report].
2. EIN Presswire. (2025) Ecer reshapes cross-border B2B trade with comprehensive AI ecosystem, driving industry-wide transformation [Press release].
3. IAB (InteractiveS Advertising Bureau) (2026) The AI ad gap widens [Industry report].
4. Song X., Hou W, Ouyang Z, Hao F (2026) AI washing: Strategic disclosure and backlash. *Finance Research Letters*, 95, Article 109684. <https://doi.org/10.1016/j.frl.2026.109684>
5. Holland & Knight (2024) Beyond the hype: The SEC's intensified focus on AI washing practices [Legal insight].
6. Magni L, Marchetti G, Alharbi A. (2023) *Learnable theory & analysis*. Luiss University Press
7. Magni L, Marchetti G, Alharbi A (2026a) A neurosemantic approach to vision and mission alignment: Construct formalization, metric design, and pilot plausibility testing. *Journal of Business Research and Reports* 3:1-22. [https://doi.org/10.47363/JBRR/2026\(3\)126](https://doi.org/10.47363/JBRR/2026(3)126)
8. Magni L, Marchetti G, Alharbi A (2024) *Learnable linguistics for business leaders*. Youcanprint
9. Magni L, Marchetti G, Alharbi A (2025) *Generative leadership of meanings*. Youcanprint.
10. Magni L (2025) Ontologies for human-centered AI: A neurosemantic framework. *Aletheia: A Journal of Literary and Linguistic Studies*, 1(23). <https://doi.org/10.5281/zenodo.17385401>
11. Magni L (2026b) When AI explains science: A hybrid-domain theory of meaning-making in human-AI science communication. Manuscript under review
12. Magni L (2026c) Hedging as Prepositional Focusing: A Learnable-Theoretic Account of the Persuasive Sweet

- Spot (July 16, 2026). Available at SSRN: <https://ssrn.com/abstract=7127282> or <http://dx.doi.org/10.2139/ssrn.7127282>.
13. Sundar S S (2008) The MAIN model: A heuristic approach to understanding technology effects on credibility. In M. J. Metzger & A. J. Flanagin (Eds.), *Digital media, youth, and credibility* 73-100 MIT Press
 14. Sundar S S, & Kim J (2019) Machine heuristic: When we trust computers more than humans with our personal information. In *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems* (pp. 1–9). ACM. <https://doi.org/10.1145/3290605.3300768>
 15. Feng Y, Sun Y (2026) Consumer empowerment through generative artificial intelligence: Enhancing brand narrative with collaboration, creation, and communication. *Review of Marketing Communications*, 45:229-261.
 16. Wen Y, Laporte S (2025) Experiential narratives in marketing: A comparison of generative AI and human content. *Journal of Public Policy & Marketing*, 44:392–410. <https://doi.org/10.1177/07439156241297973>
 17. Desveaud K., Pavone, G (2026) Consumer perceptions of AI-generated marketing content. In R. Ladhari (Ed.), *Encyclopedia of artificial intelligence in marketing*. Springer. https://doi.org/10.1007/978-3-031-75316-9_94-1
 18. Wang N, Liang C (2026) How to disclose? Strategic AI disclosure in crowdfunding (arXiv:2602.15698).
 19. Fisch J E, Laboure M., Turner I (2019) The emergence of the robo-advisor. In J. Agnew & O. S. Mitchell (Eds.), *The disruptive impact of fintech on retirement systems* (pp. 13–37). Oxford University Press.
 20. Foerster S, Linnainmaa J T, Melzer B T, Previtero A (2017). Retail financial advice: Does one size fit all? *Journal of Finance*, 72(4), 1441–1482. <https://doi.org/10.1111/jofi.12514>
 21. Ca' Zorzi M, Lopardo G, Manu A.S (2025) Verba volant, transcripta manent: What corporate earnings calls reveal about the AI stock rally (ECB Working Paper No. 3093). European Central Bank.
 22. Babina T, Fedyk A, He A. Hodson J (2024) Artificial intelligence, firm growth, and product innovation. *Journal of Financial Economics*, 151, Article 103745. <https://doi.org/10.1016/j.jfineco.2023.103745>
 23. Killman J (2023) Machine translation and legal terminology: Data-driven approaches to contextual accuracy. In H. J. Kockaert & F. Steurs (Eds.), *Handbook of terminology Legal terminology John Benjamins* 3:485-510
 24. Altakhaineh A R.M., Alghathian G A, Jarrah M. M (2025) A comparative study of accuracy in human vs. AI translation of legal documents into Arabic. *International Journal of Language & Law*, 14:63-80. <https://doi.org/10.14762/jll.2025.063>
 25. Stein J C (1989) Efficient capital markets, inefficient firms: A model of myopic corporate behavior. *The Quarterly Journal of Economics*, 104(4), 655–669. <https://doi.org/10.2307/2937861>
 26. *Macquarie Infrastructure Corp. v. Moab Partners, L.P.*, 601 U.S. 257 (2024).
 27. Federal Trade Commission (2021) *In the Matter of Everalbum, Inc.* (File No. 1923172) [Decision and order].
 28. Federal Trade Commission. (2022) *In the Matter of WW International, Inc. and Kurbo, Inc.* [Decision and order].
 29. Federal Trade Commission (2023) *In the Matter of Rite Aid Corporation* [Complaint and decision and order].
 30. Securities and Exchange Commission. (2022) SEC charges BNY Mellon Investment Adviser for misstatements and omissions concerning ESG considerations [Press release 2022-86].
 31. Global Investigations Review (2026). *Americas investigations review 2026: US enforcement agencies intensify scrutiny of AI washing*.
 32. Lin T, & Zhang, Y. (2026). Visible sources and invisible risks: Exploring the impact of AI disclosure on perceived credibility of AI-generated content. *JCOM*, 25(1), Article A09.

Copyright: ©2026 Luca Magni. This is an open-access article distributed under the terms of the Creative Commons Attribution License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.