



EUROPEAN CONFERENCE

Conference Proceedings

XXXIV International Scientific and Practical Conference
«Priority areas of scientific research in the modern
world»

August 24–26, 2026
Prague, Czech Republic

PRIORITY AREAS OF SCIENTIFIC RESEARCH IN THE MODERN WORLD

Abstracts of XXXIV International Scientific and Practical Conference

Prague, Czech Republic
(August 24–26, 2026)

9.	Ван С. КРИТЕРІЇ ВИЗНАЧЕННЯ РІВНІВ ПІДГОТОВКИ ЗДОБУВАЧІВ ВИЩОЇ ОСВІТИ КНР ДО ПРОФЕСІЙНОЇ САМОРЕАЛІЗАЦІЇ В УМОВАХ РИНКОВОЇ ЕКОНОМІКИ	36
10.	Ван С. ПЕДАГОГІЧНІ УМОВИ ПІДГОТОВКИ ЗДОБУВАЧІВ ВИЩОЇ ОСВІТИ КНР ДО ПРОФЕСІЙНОЇ САМОРЕАЛІЗАЦІЇ В УМОВАХ РИНКОВОЇ ЕКОНОМІКИ	43
11.	Дуцик А.А. СТРУКТУРНО-ФУНКЦІОНАЛЬНА МОДЕЛЬ ПІДГОТОВКИ МАЙБУТНЬОГО ВЧИТЕЛЯ ДО МЕДІАТВОРЧОСТІ У ПРОФЕСІЙНІЙ ДІЯЛЬНОСТІ	50
12.	Коврей Д.Й. ПАРТНЕРСЬКА ВЗАЄМОДІЯ ЗАКЛАДУ ДОШКІЛЬНОЇ ОСВІТИ ТА СІМ'Ї У ФОРМУВАННІ ОСОБИСТІСНОЇ КОМПЕТЕНТНОСТІ ДИТИНИ-ДОШКІЛЬНИКА	53
HISTORY AND ARCHAEOLOGY		
13.	Литвиненко О.Є. АГРОНОМІЯ В УКРАЇНСЬКИХ ГУБЕРНІЯХ РОСІЙСЬКОЇ ІМПЕРІЇ НАПРИКІНЦІ ХІХ – НА ПОЧАТКУ ХХ СТОЛІТТЯ: СТАНОВЛЕННЯ ТА РОЗВИТОК	57
JURISPRUDENCE		
14.	Veresha R. PANDEMIC CHALLENGES: KEY THREATS AND COUNTERMEASURES	61
15.	Рябенко В.П. ОСОБЛИВОСТІ ОПОДАТКУВАННЯ ІНВЕСТИЦІЙНОГО ПРИБУТКУ ЗА ЗАКОНОДАВСТВОМ УКРАЇНИ	67
MANAGEMENT		
16.	Prokopenko S.O. ORGANIZATIONAL RESILIENCE IN THE TIMES OF ALGORITHMIC DECISION SUPPORT	70
17.	Prokopenko S.O. SOVEREIGN AI AND EPISTEMIC RESILIENCE: GOVERNING NATIONAL ALGORITHMIC INFRASTRUCTURES	76

18.	Prokopenko S.O. HOW AI-GENERATED CONTENT UNDERMINES THE EPISTEMIC RESILIENCE OF BRANDS	83
19.	Prokopenko S.O. DATA SOVEREIGNTY AS A MARKETING DIFFERENTIATOR: THE RISE OF PRIVACY-FIRST POSITIONING	89
20.	Король В.С. УПРАВЛІННЯ ПРОЦЕСАМИ ФОРМУВАННЯ ЛЮДСЬКИХ РЕСУРСІВ ОРГАНІЗАЦІЇ	98
21.	Ціжма Ю.І., Ціжма О.А. ЕМОЦІЙНИЙ ІНТЕЛЕКТ ТА РЕФЛЕКСИВНЕ ЛІДЕРСТВО – ВІД ІНДИВІДУАЛЬНИХ SOFT SKILLS ДО ВИСОКОЇ ПРОДУКТИВНОСТІ ОРГАНІЗАЦІЇ	101
PHILOLOGY		
22.	Domnich V.H. LA DIVERSIDAD DIALECTAL DEL ESPAÑOL EN AMÉRICA LATINA: VARIACIÓN LINGÜÍSTICA, IDENTIDAD Y DESAFÍOS DE SU PRESERVACIÓN	104
23.	Голікова Н.С. ПАРНІ СПОЛУЧЕННЯ СЛІВ В ОПЦІЇ СУЧАСНОЇ ЛІНГВОУКРАЇНІСТИКИ: ТРАДИЦІЇ VS ІННОВАЦІЇ	108
PHILOSOPHY		
24.	Кабанець Н.І. ЕКЗИСТЕНЦІЙНЕ СПРИЙНЯТТЯ ЧАСУ В СУЧАСНОМУ ІНФОРМАЦІЙНОМУ СУСПІЛЬСТВІ	112
PSYCHOLOGY		
25.	Riabchych Y. SOCIO-PSYCHOLOGICAL DETERMINANTS OF SUICIDAL BEHAVIOR AMONG YOUNG PEOPLE	117
TECHNICAL SCIENCES		
26.	Darmofal E. AI-BASED EARLY WARNING OF TECHNOGENIC AND ENVIRONMENTAL RISKS FOR CRITICAL INFRASTRUCTURE UNDER WARTIME CONDITIONS	122

HOW AI-GENERATED CONTENT UNDERMINES THE EPISTEMIC RESILIENCE OF BRANDS

Prokopenko Serhii Olexandrovych

Doctor of Philosophy, Senior Lecturer

Simon Kuznets Kharkiv National Economic University, Department of Marketing

Abstract. The rapid diffusion of generative artificial intelligence has turned brand communication into a site of large-scale, largely unsupervised production of text, images, and videos. This article examines how AI-generated content undermines the epistemic resilience of brands as their capacity to sustain reliable, trusted, and internally consistent knowledge about themselves within a contested information environment. Drawing on recent empirical and theoretical literature, the article identifies five interacting mechanisms of erosion: the AI penalty and disclosure paradox in audience perception; authenticity discounting in influencer and social content; hallucination and the self-reinforcing "AI slop loop" in generative search; model collapse as a long-term contamination of the shared information commons; and the homogenization of brand voice under algorithmic optimization. The article argues that these mechanisms operate cumulatively rather than in isolation, converting short-term efficiency gains into long-term reputational and epistemic liabilities. It concludes with a governance framework built on provenance, human epistemic authority, and transparent disclosure practices that brands and researchers alike can use to preserve trust in AI-mediated communication.

Keywords: *brand trust; epistemic resilience; generative AI; model disclosure paradox;*

Generative artificial intelligence has moved, within a few years, from a novelty tool to the default production layer of brand communication. Product descriptions, social posts, influencer scripts, customer service replies, and even strategic reports are now routinely drafted, or entirely generated, by large language models and diffusion-based image and video systems. Industry estimates cited in recent marketing and journalism research suggest that AI-generated material already accounts for roughly half of newly published English-language web content, a share that has grown rapidly since the public release of conversational AI systems in late 2022. For brands, this shift promises speed, scale, and personalization that manual content production cannot match.

This article synthesizes recent empirical findings from marketing, communication, and AI-safety research to map the mechanisms through which AI-generated content undermines brand epistemic resilience. The analysis proceeds in three steps. First, it situates epistemic resilience as an analytical lens for brand communication in an AI-saturated environment. Second, it identifies five interacting mechanisms of erosion, moving from audience-level perception effects to systemic, infrastructure-level contamination. Third, it outlines governance principles that can

help brands and researchers preserve epistemic resilience without abandoning the efficiency gains that generative AI offers.

2. Epistemic Resilience as a Framework for Brand Communication

Epistemic resilience has been defined in institutional and political-communication research as an organization's capacity not only to discern truth amid disinformation but also to defend and communicate that truth under conditions of social upheaval and informational overload [6]. In the sustainability and organizational-strategy literature, the concept is described more broadly as the capacity of a system — an individual, an organization, or a society — to maintain or regain its ability to generate and use reliable knowledge when confronted with disturbance or uncertainty. Comparative political communication research further shows that epistemic vulnerability can be measured at the system level, and that societies and institutions with stronger public communication infrastructures tend to exhibit greater epistemic resilience than those without [3].

Transposed to brand management, epistemic resilience denotes three interlocking capacities. The first is representational accuracy: the degree to which a brand's public communication corresponds to verifiable facts about its products, policies, and claims. The second is source legitimacy: the degree to which audiences, journalists, and algorithmic intermediaries regard the brand as a credible origin of information about itself. The third is recoverability: the speed and effectiveness with which a brand can correct a false claim, whether the claim originated internally or was introduced into the information ecosystem by third parties, search engines, or AI systems. A brand with high epistemic resilience can absorb informational shocks — a viral rumor, a fabricated comparison, a hallucinated product specification — without a lasting loss of trust. A brand with low epistemic resilience finds that even a single error propagates, compounds, and becomes difficult to retract.

Generative AI intersects with all three capacities simultaneously, which is what makes its effect on brand communication systemic rather than incidental. It changes representational accuracy by introducing content that has not been authored or verified by a person with direct knowledge of the underlying facts. It changes source legitimacy because audiences and even AI systems themselves treat AI-authored and AI-mediated content differently from human-authored content. And it changes recoverability because AI-generated misinformation about a brand can be reproduced, cited, and re-cited by other generative systems faster than a communications team can issue a correction. The next section examines each of the mechanisms through which this occurs.

3. Mechanisms of Erosion

A growing body of experimental research documents what has been termed the "AI penalty"[5]: communication that is known or believed to be AI-generated is consistently rated as less trustworthy, less authentic, and less useful for forming an opinion than functionally identical human-authored communication. A large pre-registered study across workplace and social-media contexts found that this penalty affects not only perceptions of the communicator but the audience's willingness to rely on the content when deliberating [5]., which is to say that AI attribution reduces a message's epistemic uptake, independent of its factual accuracy. The same research

identified a related "disclosure paradox": audiences consistently say that disclosure of AI use is important, yet they penalize communication more heavily once that disclosure is actually made. This creates a perverse incentive structure in which honest brands that label AI-assisted content are penalized relative to brands that do not disclose, undermining the very transparency norms that would otherwise support epistemic resilience.

For brands, the practical implication is that AI-generated content does not merely risk factual errors; it risks a baseline erosion of credibility that occurs even when the content is accurate, simply because of its provenance. This dynamic is compounded when disclosure requirements — increasingly common in advertising and platform policy — collide with an audience response that punishes disclosed AI use, leaving brand communicators with no strategy that avoids some erosion of trust.

Beyond general AI attribution, research on social-media and influencer marketing shows a more specific effect: the use of generative AI to produce ostensibly personal, first-person brand content measurably diminishes perceived brand authenticity [2], even when audiences are not explicitly told that AI was involved. Experimental studies comparing AI-generated and human-created influencer content find that AI influencers reduce both perceived authenticity and brand trust relative to human influencers, and that explicit disclosure of AI involvement intensifies rather than mitigates this effect. This is particularly consequential because authenticity functions, in much of the consumer-trust literature, as the primary mediating variable between AI content exposure and downstream brand trust: content that reads as generic, formulaic, or emotionally hollow is discounted regardless of its informational accuracy [7].

A systematic review of consumer trust research on AI-generated marketing content confirms this pattern across dozens of studies, showing that perceived authenticity, together with a parallel affective reaction described as moral discomfort, consistently mediates the relationship between AI content exposure and reduced trust, moderated by factors such as consumer AI literacy, cultural context, and disclosure framing. The consistency of this finding across independent studies suggests that authenticity discounting is not a transient novelty effect but a structural feature of how audiences currently process AI-mediated brand communication [1].

A third and increasingly documented mechanism operates not at the level of audience perception but at the level of the information ecosystem itself. Generative AI systems are prone to hallucination — the confident production of fabricated facts, sources, or quotations that are indistinguishable in tone and formatting from accurate output. When such hallucinations concern a brand's pricing, features, availability, or competitive standing, they can be absorbed by other AI systems that scrape, summarize, and re-cite web content, creating what industry analysts have called an "AI slop loop": a fabricated claim about a brand is generated once, picked up and repeated by AI-driven content pipelines, cited by retrieval-augmented search and answer engines because it now has an apparent citation trail, and ultimately presented to consumers with the same confident, authoritative tone as accurate information. Because the interface and formatting of AI-generated answers is uniform regardless of accuracy, users have no practical way to distinguish a response grounded in verified

sources from one built on a fabricated claim that has simply been repeated often enough to appear well-supported.

For brand epistemic resilience, this mechanism is particularly corrosive because it removes the correction from the brand's control. Traditional misinformation could often be traced to an identifiable source — a rival, a disgruntled customer, a fake news site — and countered directly. The AI slop loop instead distributes the fabrication across many AI-mediated surfaces simultaneously, so that correcting the record with one platform or publisher does not prevent the same fabricated claim from persisting, or reappearing, on others. Brands increasingly find that a false product claim, once absorbed into generative search results, is more durable and more difficult to retract than a comparable claim published on a single conventional web page.

The fourth mechanism operates at a still deeper, infrastructural level. Foundational research published in *Nature* demonstrates that generative models trained repeatedly on data that is itself generated by earlier AI systems undergo a degenerative process known as model collapse, in which successive generations of models progressively lose the ability to represent rare, nuanced, or minority information and converge toward increasingly generic, repetitive output. Because a substantial and growing share of newly published web content is now itself AI-generated, the training data available to future AI systems — including the systems that brands rely on for search visibility, customer-facing chatbots, and generative-engine optimization — is increasingly self-referential rather than grounded in independently verified, human-authored sources.

For brand communication, the practical consequence of model collapse is a slow but systemic degradation of the shared information commons in which brand claims circulate. As AI systems train on progressively more AI-generated material, distinctive, brand-specific information is statistically the kind of "rare" or "tail" detail that model collapse research shows is lost first, while generic, repeated claims about a category or competitor are reinforced. Over time, this can flatten legitimate brand differentiation in the very AI systems that consumers increasingly consult before purchase, replacing accurate but uncommon details with the statistically dominant, and not necessarily correct, consensus circulating across AI-generated sources.

A final, more subtle mechanism concerns not factual accuracy but distinctiveness. Because generative language models are trained to produce the statistically most probable continuation of a prompt, large-scale reliance on such systems for brand copywriting tends to converge different brands toward similar vocabulary, sentence rhythm, and rhetorical structure. Comparative studies of AI-generated marketing narratives, reviews, and social content find that such material is more textually similar to other AI-generated material than to the human-authored content it is meant to resemble, indicating a homogenizing effect that operates independently of factual accuracy.

This homogenization undermines epistemic resilience indirectly but persistently: a brand's distinctive voice has historically functioned as a heuristic signal of authenticity and accountability, allowing audiences to attribute claims to a specific, identifiable source with a track record. As AI-generated brand content converges

stylistically across competitors, this heuristic weakens, making it harder for audiences to distinguish a brand's genuine communication from that of competitors, imitators, or, in the case of parody and fraud, malicious impersonators.

4. Consequences for Brand Epistemic Resilience

Taken together, these mechanisms compound rather than operate independently. A brand that adopts generative AI for content production first experiences a baseline discount in perceived authenticity and trust, described above as the AI penalty. If a disclosed claim later proves inaccurate — because of hallucination — the brand experiences a second, sharper credibility loss, this time compounded by the disclosure paradox, since honest brands that had disclosed AI involvement are penalized more heavily than brands that concealed it. If the inaccurate claim propagates through the AI slop loop before it can be corrected, the brand loses control over the timeline and channels of correction. Over a longer horizon, as model collapse degrades the training data available to future AI systems, the brand's accurate, distinctive claims become statistically harder for AI systems to retrieve and represent correctly, while homogenized, generic brand voice makes it harder for audiences to recognize genuine communication when it does appear.

The cumulative effect is a brand that is simultaneously more prolific and less epistemically trustworthy — publishing more content, across more channels, while its capacity to be believed, cited accurately, and distinguished from competitors or impersonators steadily declines. This is the central paradox this article identifies: the same technology that expands a brand's capacity to communicate can, without deliberate governance, shrink its capacity to be believed.

Addressing this paradox requires governance measures at three levels. At the content level, brands should maintain a human-verified factual layer — a single, authoritative, and clearly dated source of product specifications, pricing, and claims — that both human editors and AI systems can be pointed to, reducing the surface area available for hallucination and stale or fabricated claims to spread. At the disclosure level, rather than treating AI disclosure as a binary label that triggers the disclosure paradox, brands and researchers should explore framing strategies documented in recent trust research, such as pairing disclosure with evidence of human oversight and accountability, which experimental studies suggest can attenuate, though not eliminate, the AI penalty.

At the infrastructural level, brands have a growing incentive to participate in content-provenance initiatives — cryptographic and metadata-based standards that allow AI systems and platforms to distinguish verified, brand-authored content from third-party or synthetic material — and to monitor how major generative search and answer engines represent their brand, treating this monitoring as a core reputational function rather than a peripheral technical concern [6]. Finally, at the research and policy level, the emerging norm in academic publishing of requiring explicit, itemized disclosure of AI tool use, together with an assignment of human responsibility for accuracy, offers a transferable model: just as scholarly journals now require author teams to take personal responsibility for the accuracy of AI-assisted research, brands can adopt an analogous internal standard in which a named human owner is

accountable for the factual accuracy of any AI-assisted public claim, regardless of how much of its drafting was automated.

6. Conclusion

Generative AI has not created a single new threat to brand communication so much as it has activated and accelerated several pre-existing vulnerabilities in how trust, authenticity, and factual accuracy are produced and maintained. The AI penalty and disclosure paradox show that audiences discount AI-mediated communication even when it is accurate; authenticity discounting shows that this discount is sharper still for personal or influencer-style content; the AI slop loop shows how a single fabricated claim can be laundered into an apparent fact across generative search systems. None of these mechanisms is fully solved by simply using AI more carefully or more sparingly.

References

1. Kirill Baryshkov, Yana Kuzina, Iryna Smuk, Mila Tkachuk (2026). Consumer Trust in AI-Generated Marketing Content: A Systematic Literature Review and Research Agenda. *American Impact Review*. <https://doi.org/10.66308/air.e2026024>
2. Brüns, J. D., & Meißner, M. (2024). Do you create your content yourself? Using generative artificial intelligence for social media content creation diminishes perceived brand authenticity. *Journal of Retailing and Consumer Services*, 79, Article 103790. <https://doi.org/10.1016/j.jretconser.2024.103790>
3. Labarre, J. (2025). Epistemic Vulnerability: Theory and Measurement at the System Level. *Political Communication*, 42(1), 6–26. <https://doi.org/10.1080/10584609.2024.2363545>
4. Anthropic. (2026). *Claude Sonnet 5* (Aug. 2026 version) [Large language model]. <https://claude.ai>
5. Sahebi, S., Formosa, P., & Bankins, S. (2026). The AI penalty and disclosure paradox: Trust, authenticity and knowledge uptake in AI-mediated communication. *Computers in Human Behavior: Artificial Humans*, 8, Article 100304. <https://doi.org/10.1016/j.chbah.2026.100304>
6. The Decision Lab. (2026). Engineering epistemic resilience in information ecosystems. <https://thedecisionlab.com/big-problems/engineering-epistemic-resilience-in-information-ecosystems>
7. Ujjainwala, F. J. (2025). Influence of AI-generated influencer content on brand trust and authenticity perceptions. *Journal of Marketing & Social Research*, 2(9), 256–262.

Scientific publications

MATERIALS

The XXXIV International Scientific and Practical Conference
«Priority areas of scientific research in the modern world»

Prague, Czech Republic
(August 24–26, 2026)