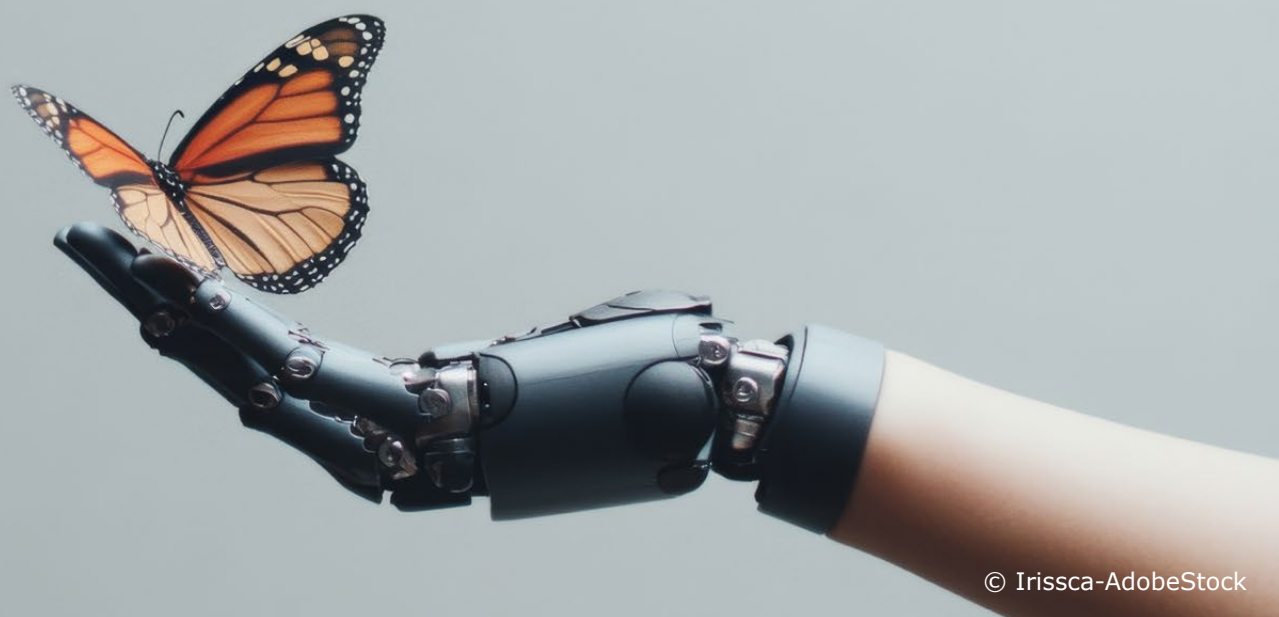




© Irissca-AdobeStock

LIGHTS - Human/Well being

Fast answers, shallow thinking: reclaiming reflexivity in an accelerated AI world

ESCP Impact Paper No.2025-23-EN

Diana GARCIA QUEVEDO, Anna GLASER & Caroline VERZAT
ESCP Business School

[Fast answers, shallow thinking: reclaiming reflexivity in an accelerated AI world

Diana Garcia Quevedo*

Anna Glaser**

Caroline Verzat**

ESCP Business School

Abstract

As organizations increasingly rely on generative AI tools, there is a growing risk of mistaking speed and surface plausibility for understanding. Building on Messeri & Crockett (2024) and our own research (Garcia Quevedo, Glaser, & Verzat, 2025), this impact paper examines four key “illusions of understanding” that AI can foster (illusions of depth, breadth, objectivity, and relevance) and which distort how knowledge is perceived, produced, and acted upon. While Messeri & Crockett identify the cognitive risks of over-trusting AI outputs, our research highlights how these risks manifest in qualitative analysis and theorization, particularly when interpretive processes are not reflexively or collectively maintained. To address this, we draw on Cunliffe’s (2004, 2016) concept of reflexivity as a relational, ethical, and situated practice of questioning assumptions and examining how knowledge is constructed. We propose reflexivity and knowledge diversity as safeguards for responsible AI use. In a world of fast answers, more reflexive practices might invite us to slow down and think more deeply together.

Keywords: artificial intelligence (AI), reflexivity, illusions of understanding, knowledge diversity, acceleration

*PhD graduate, ESCP Business School

**Professor, ESCP Business School

ESCP Impact Papers are in draft form. This paper is circulated for the purposes of comment and discussion only. Hence, it does not preclude simultaneous or subsequent publication elsewhere. ESCP Impact Papers are not refereed. The form and content of papers are the responsibility of individual authors. ESCP Business School does not bear any responsibility for the views expressed in the articles. Copyright for the paper is held by the individual authors.

Fast answers, shallow thinking: reclaiming reflexivity in an accelerated AI world

Introduction

Organizations are accelerating the introduction of artificial intelligence (AI) in every aspect of their workflow. In May 2025, *Les Echos* reported that Duolingo's CEO declared that the company was fully integrating AI, even restricting further hiring to areas where automation is not possible (Boone, 2025). In academia, a similar pattern has been observed, with researchers adopting AI in every aspect of their work. AI tools offer automation capabilities that enable scaling and allow organizations to better serve their existing offerings and even expand them. As such, AI has become more than a managerial ambition, it has marked a radical shift in organizational culture.

At the same time, the risks associated with the limitations and ethical challenges of AI systems are becoming increasingly evident. *The New York Times* reports that the latest generation of AI “reasoning systems” produces factual errors, so-called “hallucinations”, at rates as high as 79% (Metz & Weise, 2025). Furthermore, *The Guardian* (Hall, 2025) has uncovered AI-authored books on sensitive medical topics, such as ADHD, filled with “dangerous nonsense” and sold without disclosure or expert oversight. More recently, a philosophical essay introducing the concept of “hypnocracy” was published under the name of Jianwei Xun, a supposedly Hong Kong-based thinker, who turned out to be a fictional AI-generated persona (Riché, 2025). The case sparked intense debate about the boundaries of authorship and academic publishing.

On one hand, AI tools, particularly generative ones, have important limitations linked to their probabilistic natures, the outputs of which can be erroneous but still sound plausible. On the other hand, users may perceive outputs as unbiased and factual, paving the way to possible misuse and fraudulent behavior. Together, these limitations and challenges paint a troubling picture: AI systems are being integrated into core service and knowledge functions, even if reliability is a concern.

This Impact Paper examines the pressing question: How can AI systems be effectively integrated into organizational workflows when lack of reliability is apparent? How can we think responsibly when the pressure to adopt AI is so high? Building on recent research, we argue that current uses of AI often lead to reliance on outputs that appear convincing rather than truly understood, and to decision-making that favors speed and efficiency over critical reflection. To address this, we call for a renewed focus on reflexivity (Cunliffe, 2004, 2016), understood as a way of questioning our assumptions and engaging with others in context-aware, relational, and morally responsive ways. These elements are always important, but particularly when working with AI.

The illusion of understanding in the age of acceleration

The integration of AI into organizational workflows is expanding rapidly, with tools increasingly employed from the initial conception of projects and services to the final public deployment. In academic research, for instance, large language models (LLMs) are now used to generate interview questions, summarize articles, and even draft full publications. Their appeal lies in the promise of speed, productivity, and apparent precision. However, this growing reliance raises a pressing concern: what kinds of knowledge, products, and services are being produced, and what kinds are being lost?

Messeri & Crockett (2024) caution against what they call illusions of understanding, cognitive traps that emerge when users take AI outputs at face value. These illusions reflect deeper cognitive risks: the dangers of forming or acting upon false beliefs about what we know, especially when AI systems create a misleading impression of depth, breadth, neutrality, or relevance. These risks manifest in four main ways: the illusion of explanatory depth, whereby users believe they understand more than they do; the illusion of exploratory breadth, whereby algorithmically constrained suggestions are mistaken for comprehensive knowledge; the illusion of objectivity, whereby AI-generated results are wrongly perceived as neutral and bias-free; and the illusion of relevance, whereby users neglect critical and ethical thinking when applying, evaluating or creating AI-generated knowledge. When these illusions take hold, critical interpretation is replaced by superficial plausibility, contextualized ethical thinking is replaced by disembodied legitimacy, and knowledge becomes automated rather than reflective.

In his blog, "The 70% problem: Hard truths about AI-assisted coding," Osmani (2024), a lead engineer at Google, clearly pictures these illusions of understanding. AI tools initially enabled software engineers to work faster, but without critically questioning and revising AI outputs, this initial advantage was lost. Eventually, the reliance on an AI-assisted tool resulted in diminishing returns and poor-quality products.

This view is echoed by Roberts et al. (2024), who argue that over-reliance on AI risks eroding the skills of qualitative researchers. While AI can speed up many research tasks, it also projects an illusion of human-like capability. These tools may generate text that appears insightful, but they lack genuine understanding, ethical sensibility, and the capacity to acknowledge uncertainty. Left unchecked, such illusions could lead to the deskilling of researchers and the marginalization of interpretive labor.

These risks are amplified in environments marked by what Rosa (2003) has termed social acceleration, a condition where the pressure to do more, faster, reshapes not only how we act but how we think. AI is often positioned as a remedy to time scarcity, promising to reduce cognitive load and increase efficiency. Yet as Rosa (2003) shows, acceleration tends to diminish our capacity for deliberation, deepens our dependency on external aids, and ultimately weakens our relationship to the world. In this context, AI becomes not a neutral tool but a driver of contraction: the faster we go, the less we engage.

The problem extends beyond individuals to the structure of knowledge itself. Messeri & Crockett (2024) warn that overuse of AI may foster scientific monocultures: narrow norms of inquiry that emerge when research questions are shaped by what AI can handle, rather than what is most important or relevant. These monocultures arise in two forms: monocultures of knowing, which privilege methods suited to automation, and monocultures of knowers, which devalue cognitive and demographic diversity by assuming AI's universality.

Taken together, these concerns highlight a systemic risk: that AI, if left unexamined, will reshape practices around speed, surface plausibility, and disembodied legitimacy. In doing so, it risks hollowing out the very processes of interpretation and reflection that make value creation meaningful. If we are to use AI responsibly, we must remain vigilant, not only about what it can do, but about how it changes what we do when we rely on it too much.

Reclaiming reflexivity: toward situated, diverse, ethical interpretation

If illusions of understanding and acceleration-driven shortcuts pose cognitive risks (Messerli & Crockett, 2024), how can we restore depth and responsibility in AI-supported work? One critical path forward is to reclaim reflexivity, not as a vague call for caution, but as a situated, relational, and ethical practice of knowledge-making.

Cunliffe (2004, 2016) defines critical reflexivity as the process of questioning the assumptions, values, and actions that shape how we think and act in organizations. It requires us to see knowledge not as objective data detached from context, but as something created in relation to others, with systems, and with the histories and identities we carry. Reflexivity is thus not just a personal skill; it is a collective practice of relating differently to knowledge, recognizing how we co-construct meaning, how our worldviews shape what we notice or ignore, and how ethical action emerges from this awareness.

In AI-mediated environments, such reflexivity is urgently needed. As tools like LLMs become embedded in everyday decision-making, the risk is not only that we outsource judgment, but also that we begin to forget what judgment requires: the ability to navigate uncertainty, acknowledging disagreement, and taking context seriously. Reflexivity encourages us to question what we take for granted (Cunliffe, 2004, 2016), including how knowledge is produced and whose perspectives it reflects (Stinson & Vlaad, 2024; Browne et al., 2024). This includes asking questions such as: “Is this result useful?”, “How was it produced?”, “Whose perspective is embedded in it?”, and “What might it leave out?”.

Collaborating with AI tools starts with reflecting on AI capabilities and distinctive human skills. At four of the six knowledge levels of Bloom’s taxonomy (understand, apply, evaluate, and create), AI can provide expanded data, alternatives, predictions, rubrics, etc (Ecampus Oregon State University, 2024). However, human beings play an essential role in critically reviewing AI output. Therefore, now all learners, teachers, researchers, or decision makers should consciously design and account for their knowledge processes step-by-step, clarifying what is devoted to AI and human beings. Moreover, they should use precise criteria to assess the cognitive validity of AI output (clarity, accuracy, precision, breadth, depth, relevance, and significance) and metacognitive criteria expected from humans (holistics, ethics, authenticity, originality) in each specific situation (Zaphir et al., 2024).

This call for reflexivity is not a retreat into individual introspection. Rather, as Stinson & Vlaad (2024) argue, robust interpretation is an emergent property of diverse, dialogic teams. They introduce the concept of emergent expertise, the idea that no single person, however skilled, can fully grasp complex problems alone. Instead, meaningful insight arises through the interaction of different kinds of expertise, technical, experiential, and disciplinary, especially when those differences are allowed to surface in inclusive ways.

Rostcheck & Scheibling (2024) further develop the importance of incorporating diverse teams in AI systems. They argue that a single vision of expertise cannot possibly capture the vast possible interactions between humans and AI tools. Diverse teams, including psychologists, social workers, negotiators, and others with expertise in human interaction, can bring a broader range of perspectives and strategies to better align AI systems with complex human realities. They draw a parallel with the fable of blind researchers examining an elephant, where each researcher’s limited perspective necessitates integrating information from diverse viewpoints for a comprehensive understanding. Similarly, AI systems require integrating insights from various disciplines to navigate the complex and rapidly evolving landscape effectively (Rostcheck & Scheibling, 2024). As such, diverse teams suggest a more robust and adaptable approach to ensuring AI safety. In this view, diversity

is not an ideal but a cognitive necessity: it expands the range of what can be noticed, questioned, and understood.

In our own research on LLMs in qualitative analysis (Garcia Quevedo et al. 2025), we propose a method that reflects this commitment to reflexive pluralism. While we utilize AI tools to assist researchers in navigating large and complex datasets, we emphasize that the interpretive core must remain collective and human led. Our approach does not treat AI as a neutral assistant, but as a tool whose limitations must be acknowledged and whose outputs require active, situated interpretation. We recommend combining different models, triangulating results, and, crucially, bringing together researchers with diverse backgrounds to interrogate both the data and the tools themselves.

Taken together, these perspectives offer a clear alternative to the illusions of understanding described earlier. Rather than seeing reflexivity and diversity as soft or optional, we position them as structural conditions for critically evaluating and validating AI systems. Without reflexive and diverse interpretive practices, organizations risk treating AI outputs that sound plausible as if they reflected reliable knowledge, mistaking fluent or consistent language for deep understanding, and broad agreement around AI-generated results for actual validity or critical scrutiny.

Calling for reflexivity means valuing how we produce knowledge, how we build diverse teams, and how we evaluate AI outputs. It asks us to value discomfort, dialogue, and slowness, not as inefficiencies, but as signs that we are doing the hard work of thinking together. In this way, reflexivity becomes not just a safeguard against AI's risks and limitations but a foundation for more thoughtful, inclusive, and accountable knowledge practices.

Reflexive incorporation of AI in organizations

As AI becomes more deeply embedded in organizational processes, the risks of overconfidence, cognitive shortcuts, and the loss of critical interpretation are no longer abstract concerns. Our impact paper shows that reflexivity and knowledge diversity are not just desirable qualities; they are essential safeguards for the responsible and meaningful use of AI.

To navigate these challenges, organizations must shift their focus from “what AI can do” to “how we make sense of what it does”, and who is involved in that sensemaking. This means embedding reflexive practices across the AI lifecycle: from initial problem framing and proof of concept to evaluation and validation of results.

We propose four concrete ideas to support this shift:

- *Reflexive training* for both technical staff and users, aimed at developing their capacity to critically reflect on the collaborative knowledge-building process between humans and AI. This includes questioning taken-for-granted assumptions, making explicit the values that shape decisions, and engaging ethically with AI outputs (for inspiration, see Cunliffe (2016), who emphasizes reflexivity as a relational, ethical, and situated practice of questioning what we take for granted).
- *Inclusive development processes* that involve people with different roles and experiences - including those most affected by how AI is used - to ensure systems reflect diverse needs (for inspiration, see Browne et al. (2024) who show that effective AI ethics must be grounded in lived perspectives and inclusive design, not detached principles)

- *Team-based evaluation frameworks*, such as interdisciplinary teams that bring together technical, operational, and contextual expertise to evaluate AI systems in practice (for inspiration, see Stinson & Vlaad (2024) who argue that robust interpretation emerges from diverse teams, where epistemic insight is co-constructed)
- *Time and space for interpretation*, not just automation. Organizations should protect deliberative time to avoid premature decisions based on AI outputs, especially under performance pressure (for inspiration, see Rosa (2003), who warns that social acceleration undermines our ability to reflect, making time a condition for ethical and informed action).

IDEO, a leading company renowned for its human-centered approach to innovation, has incorporated some of the previous ideas. By building on its human-centered approach, IDEO recently introduced a powerful, practical tool for fostering the ethical implementation of AI: the AI Ethics Cards (IDEO, 2025). These cards offer concrete prompts and collaborative exercises that support teams in anticipating blind spots, challenging assumptions, and re-centering on human needs throughout the AI implementation process. Rather than offering prescriptive answers, they encourage situated, collective reflection, inviting employees to ask, What are we not seeing? Whose voices are missing? What unintended consequences could arise? By embedding structured moments of reflection into AI use, IDEO helps teams resist the illusion of understanding and remain attentive to complexity.

All these recommendations align closely with ESCP's commitment to leadership, human well-being, and responsible technology. Reflexivity complements human-centered leadership, enriches ethical business practice, and supports inclusive innovation by expanding the scope of who counts as a knower in AI-driven systems.

Ultimately, this paper calls for a rebalancing of AI enthusiasm with interpretive humility. If we are to make AI work for organizations and society, not just efficiently but ethically, we must protect space for transparency, dissent, dialogue, and deliberation. These are not barriers to innovation, they are its foundation. In a world of fast answers, only reflexivity protects us from shallow thinking.

References

- Boone, J. (2025), "C'est la meilleure décision que nous avons prise récemment" : Duolingo dit adieu à ses prestataires et les remplace par l'IA. May 2nd, *Les Echos*.
<https://www.lesechos.fr/travailler-mieux/vie-au-travail/cest-la-meilleure-decision-que-nous-avons-prise-recemment-duolingo-dit-adieu-a-ses-prestataires-et-les-remplace-par-lia-2163194>
- Browne, J., Drage, E., & McInerney, K. (2024), Tech workers' perspectives on ethical issues in AI development: Foregrounding feminist approaches. *Big Data & Society*, 11(1)
- Cunliffe, A. L. (2004), On Becoming a Critically Reflexive Practitioner. *Journal of Management Education*, 28(4), 407-426.
- Cunliffe, A. L. (2016), "On becoming a critically reflexive practitioner" redux: What does it mean to be reflexive?. *Journal of Management Education*, 40(6), 740-746.
- Ecampus Oregon State University (2024). Bloom's Taxonomy Revisited – Artificial Intelligence Tools – Faculty Support, Oregon State Ecampus.
<https://ecampus.oregonstate.edu/faculty/artificial-intelligence-tools/blooms-taxonomy-revisited/>

Garcia Quevedo, D., Glaser, A., & Verzat, C. (2025), Enhancing Theorization Using Artificial Intelligence: Leveraging Large Language Models for Qualitative Analysis of Online Data. *Organizational Research Methods* (forthcoming)

Hall, R. (2025), 'Dangerous nonsense': AI-authored books about ADHD for sale on Amazon. May 4th, *The Guardian*.

<https://www.theguardian.com/technology/2025/may/04/dangerous-nonsense-ai-authored-books-about-adhd-for-sale-on-amazon>

IDEO (2025), AI Needs an Ethical Compass. This Tool Can Help.

<https://www.ideo.com/journal/ai-needs-an-ethical-compass-this-tool-can-help>

Messeri, L., & Crockett, M. J. (2024), Artificial intelligence and illusions of understanding in scientific research. *Nature*, 627(8002), 49-58.

Metz, C., & Weise, K. (2025), A.I. Is Getting More Powerful, but Its Hallucinations Are Getting Worse. May 5th, *The New York Times*.

<https://www.nytimes.com/2025/05/05/technology/ai-hallucinations-chatgpt-google.html>

Osmani, A. (2024), "The 70% Problem: Hard Truths about AI-Assisted Coding." Substack newsletter. December 4th, *Elevate* (blog).

<https://addyo.substack.com/p/the-70-problem-hard-truths-about>.

Riché, P. (2025), "Le philosophe Jianwei Xun n'existe pas, mais son concept d'« hypnocratie » entre dans notre réalité". April 16th, *Le Monde*.

https://www.lemonde.fr/idees/article/2025/04/16/le-philosophe-jianwei-xun-n-existe-pas-mais-son-concept-d-hypnocratie-entre-dans-notre-realite_6596657_3232.html

Roberts, J., Baker, M., & Andrew, J. (2024), Artificial intelligence and qualitative research: The promise and perils of large language model (LLM) 'assistance'. *Critical Perspectives on Accounting*, 99, 102722.

Rosa, H. (2003), Social acceleration: ethical and political consequences of a desynchronized high-speed society. *Constellations: An International Journal of Critical & Democratic Theory*, 10(1).

Rostcheck, D., & Scheibling, L. (2024), "The Elephant in the Room -- Why AI Safety Demands Diverse Teams." *arXiv*.

<https://doi.org/10.48550/arXiv.2407.10254>.

Stinson, C., & Vlaad, S. (2024), A feeling for the algorithm: Diversity, expertise, and artificial intelligence. *Big Data & Society*, 11(1)

Zaphir, L., Lodge, J. M., Lisec, J., McGrath, D., & Khosravi, H. (2024), "How critically can an AI think? A framework for evaluating the quality of thinking of generative artificial intelligence". *arXiv preprint arXiv:2406.14769*