



## Does Expertise Still Matter? Generative AI and the Crisis of Professional Identity

Online First

Dr. Lucy Michael Nyagoga<sup>§\*</sup>

ORCID 0000-0001-6291-5159

\*Corresponding Author



§ School of Educational Technology, Faculty of Education, Southwest University, Chongqing, China

RECEIVED

2026-06-30

ACCEPTED

2026-07-10

ONLINE PUBLISHED

2026-08-03

PEER REVIEW

Double Blind

### Abstract

The rapid diffusion of generative artificial intelligence (GenAI) tools has opened unimagined avenues for disrupting higher education, enabling professionals, especially researchers, to produce expert-seeming outputs and claim “expert-level status” without formal training in artificial intelligence. Yet, despite its popularity, concerns about AI’s effects on labor and expertise have largely overlooked a deeper categorical crisis: the collapse of the distinction between AI tool proficiency and genuine AI expertise and the consequences this has for professional identity and institutional decision-making. This positional paper interrogates the boundary between AI use and AI expertise, arguing that access to a tool that simulates expert output creates conditions under which the distinction between literacy and expertise becomes nearly impossible to perceive. Drawing on AI literacy frameworks, expertise theory, and professional identity. The paper maps this crisis by introducing simulated contributory expertise as a construct to name this condition at both individual and collective levels. The paper raises pressing implications for how higher education certifies expertise, designs AI literacy curricula, and governs institutional decision-making in an era where the professional identity, AI expertise, and its substance have become difficult to pinpoint.

Artificial intelligence (AI)

Generative artificial intelligence (GenAI)

Professional identity

AI literacy

higher education

### AI USE STATEMENT

No generative AI was used for analysis or results.

### FUNDING

No external funding was declared for this work.

### CONFLICT OF INTEREST

The authors declare no conflict of interest.

### DATA AVAILABILITY

Not applicable for this article.

### ETHICS

No ethics committee approval was required for this article type.

### CONSENT

Not applicable for this article.

### TRIAL REG.

Not applicable.

Crossref DOI: 10.34257/GJHSSG258994

**How to Cite:** Michael Nyagoga (2026). Does Expertise Still Matter? Generative AI and the Crisis of Professional Identity. Global Journal of Human-Social Science, Ahead of Print, 1-15. DOI: 10.34257/GJHSSG258994

### LICENSE

© 2026 Global Journals. Open-access article under CC BY-NC-ND 4.0 International License.



Print ISSN 0975-587X



9 770975 587004

Online ISSN 2249-460X



9 772249 460013

Under the strict compliance and defined process of



METADATA CONTINUATION

AUTHOR CONTACT QR LEDGER

Dr. Lucy Michael Nyagoga\$\*



ARCHIVAL RECORD

# Does Expertise Still Matter? Generative AI and the Crisis of Professional Identity

Dr. Lucy Michael Nyagoga<sup>§\*</sup> 

## Affiliations

<sup>§</sup> School of Educational Technology, Faculty of Education, Southwest University, Chongqing, China

## Abstract

The rapid diffusion of generative artificial intelligence (GenAI) tools has opened unimagined avenues for disrupting higher education, enabling professionals, especially researchers, to produce expert-seeming outputs and claim “expert-level status” without formal training in artificial intelligence. Yet, despite its popularity, concerns about AI’s effects on labor and expertise have largely overlooked a deeper categorical crisis: the collapse of the distinction between AI tool proficiency and genuine AI expertise and the consequences this has for professional identity and institutional decision-making. This positional paper interrogates the boundary between AI use and AI expertise, arguing that access to a tool that simulates expert output creates conditions under which the distinction between literacy and expertise becomes nearly impossible to perceive. Drawing on AI literacy frameworks, expertise theory, and professional identity. The paper maps this crisis by introducing simulated contributory expertise as a construct to name this condition at both individual and collective levels. The paper raises pressing implications for how higher education certifies expertise, designs AI literacy curricula, and governs institutional decision-making in an era where the professional identity, AI expertise, and its substance have become difficult to pinpoint.

**Keywords:** *Artificial intelligence (AI), Generative artificial intelligence (GenAI), Professional identity, AI literacy, higher education*

## \* Corresponding Author

Dr. Lucy Michael Nyagoga

## DOI

10.34257/GJHSSG258994

## 1. Introduction

In early 2023, a senior academic administrator at a research-intensive university declared during a faculty meeting that their institution no longer needed to hire “expensive AI consultants” because “everyone has ChatGPT now.” The statement was met with murmurs of agreement from some quarters and visible discomfort from others, particularly from the computer science faculty who had spent decades building the expertise now apparently rendered obsolete by a tool anyone could access for twenty dollars a month. This moment, unremarkable in its particulars, because it has been repeated in countless variations across universities, hospitals, law firms, and corporations, captures the phenomenon this paper seeks to name and interrogate.

The diffusion of generative artificial intelligence (GenAI) tools has occurred at an unprecedented pace in the history of technological adoption. Within eighteen months of ChatGPT’s public release, knowledge workers across sectors had integrated large language models into their daily workflows (Dwivedi et al., 2023). This diffusion has produced a curious and underexamined phenomenon: individuals with no formal training in artificial intelligence now routinely describe themselves, and are described by others, as “AI experts” based solely on their proficiency in using AI but not necessarily classified as AI experts. The academic administrator who believes that their institution’s AI needs are solved by ChatGPT access is not lying. They genuinely believe that they understand something that they do not necessarily understand. This belief, multiplied across thousands of institutions and millions of knowl-

edge workers, constitutes something more consequential than individual error. It represents an emerging category crisis in what it means to be an AI-related expert.

This position paper comprises four sections. First, I establish a conceptual distinction between possessing AI literacy and being an AI expert, drawing on established frameworks that reveal why functional proficiency cannot substitute for genuine understanding. Second, I argue that GenAI introduces a distinctive form of identity threat that existing identity work scholarship has not adequately theorized. Third, I draw insights the central paradox from the few empirical studies that Gen AI compresses the measurable expertise gap while creating conditions under which users cannot recognize when their expertise is failing. Finally, I articulate the implications for higher education and professional practice and offer tentative solutions, arguing that unless we develop a more precise vocabulary for distinguishing use from understanding, we will continue to build consequential decisions on the foundations of simulated expertise.

## 2. “AI Literacy” Versus “AI Expert”

The AI literacy literature provides a necessary starting point for distinguishing tool use from genuine understanding. Long and Magerko (2020) defined AI literacy as a set of competencies that enables individuals to critically evaluate, communicate with, and effectively use AI technologies. Their framework encompasses knowing what AI is, understanding how it works, recognizing when and where it should be deployed, and grasping its ethical

implications. Critically, operating an AI interface (i.e., typing a prompt into ChatGPT or any other AI tool), fulfills only a fraction of these competencies. A user who can obtain useful outputs from GenAI has achieved what Ng et al. (2021) term "functional use", but not necessarily achieved "*conceptual understanding*." This distinction is not pedantic. It is the difference between knowing that a prompt produces a certain output and knowing *why* the output is structured as it is, what the model's architecture enables and constrains, where the training data originated from, and under what conditions the output is likely to be unreliable.

This distinction becomes sharper when examined through the lens of expertise theory. Dreyfus and Dreyfus (1986) proposed that expertise develops through five stages (novice, advanced beginner, competent, proficient, and expert) and that genuine expertise *cannot* be reduced to rule-following or tool operation. Novices follow context-free rules. Experts perceive situations holistically and respond with tacit, embodied judgments that cannot be fully articulated. When a novice uses GenAI to produce expert-seeming output, they have not traversed these five stages. They have borrowed or acquired the appearance of their destination without undertaking the journey. The output may be indistinguishable from expert work, but the capacity that matters most (knowing when the output is wrong, understanding why it is wrong, and deciding what to do about it) never develops, because it is never exercised. The hallmarks of genuine expertise therefore remain absent even when the product looks expert.

To paint a distinctive picture, Collins and Evans (2007) offer further precision by distinguishing between "*contributory expertise*" (the ability to perform a practice) and "*interactional expertise*" (the ability to talk meaningfully about a practice without being able to perform it). GenAI users occupy an ambiguous position relative to this typology. They can produce outputs that resemble like the work of a contributory expert, yet lack the embodied judgment that undergirds genuine contributory expertise. They can talk about the outputs - "I asked ChatGPT to analyze this dataset and it identified three key trends" - but this talk is not interactional expertise in Collins and Evans (2007) sense, because it does not demonstrate deep understanding of the specific field's context. Rather, something I call *simulated contributory expertise*: the ability to generate the outputs of a practice without the embodied judgment that ordinarily produces them. The simulation is not fraudulent. The user genuinely believes they understand because the tool's fluency creates what Bender et al. (2021) term an "illusion of meaning." The statistical pattern-matching that produces fluent text is mistaken for comprehension, both in the machine and the user.

One might object that this is nothing new. Calculators, spell-checkers, and statistical packages such as SPSS have long let users produce competent outputs without the underlying mathematical or statistical knowledge - a form of cognitive offloading that predates GenAI by decades (Risko & Gilbert, 2016). The objection deserves a direct answer, and the answer is that it misidentifies what is offloaded. Earlier tools automated bounded, rule-governed sub-tasks within closed domains: a calculator returns a determinate answer the user can in principle verify, and SPSS executes a procedure the user has already selected. The interpretive layer (deciding which test to run, judging whether a result is plausible, determining what it means) stayed with the user.

GenAI differs in three respects. First, it does not automate a sub-task but simulates the judgment-bearing output itself: the argument, the synthesis, the analysis that constitutes the practice. Second, it operates in open-ended domains where no determinate

correct answer exists for the non-expert to check against. Third, it delivers that output in fluent, confident prose that imitates the surface signature of understanding, so the very cue that once flagged incompetence - inability to produce expert-looking work - is removed. The novelty is not offloading. It is that the offloaded layer is the one that constitutes expertise, and that the tool conceals the offload from the user.

To further highlight why AI literacy and AI experts are not the same, the Dunning-Kruger effect (Kruger & Dunning, 1999) provides the final piece of this theoretical architecture. Individuals with the least competence in a specific field's context are most likely to overestimate their expertise. GenAI amplifies this effect by providing fluent and confident outputs that mask the user's underlying lack of understanding. The tool makes the user feel competent because its outputs appear competent. The metacognitive deficit that would ordinarily signal incompetence (the inability to produce work that meets expert standards) is short-circuited by the tool's capacity to produce such work. The user receives positive feedback on their GenAI-augmented outputs and reasonably concludes that they are developing genuine expertise in the process. They are not. They are developing proficiency in a tool whose outputs they cannot evaluate reliably.

To concur, Floridi and Chiriatti (2020) put the matter starkly: GPT-3 produces text without understanding. Its outputs are syntactically impressive but semantically hollow. Confusing its output with intelligence is a categorical error. I extend this insight by arguing that confusing one's ability to produce such outputs with one's own AI expertise is the same category error, displaced onto the self. Users attribute to themselves a competence that properly belongs to the tool, and more precisely, to the tool's developers and the vast corpus of human-generated text on which it was trained.

### 3. The Professional Identity Crisis Nobody Is Talking About

Professional identity is not a stable attribute but rather a discursive accomplishment. Alvesson and Willmott (2002) established that individuals construct and maintain professional selves through ongoing identity work from formation, repair, and revision of self-narratives in organizational life (Svenningsson & Alvesson, 2003). This identity work is constrained by organizational identity regulation, the discourses and practices through which institutions shape what counts as a legitimate professional self. For academics, this regulation operates through peer review, promotion criteria, citation metrics, and the informal prestige economy of a disciplinary reputation. For knowledge workers more broadly, it operates through job descriptions, performance evaluations, and tacit recognition of who "knows and does things" and who does not.

Technological disruptions have long been recognized as triggers of intensified identity work. Ibarra (1999) demonstrated that professionals construct "provisional selves" during role transitions, experimenting with possible identities before committing to one. Pratt et al. (2006) showed that professional identity construction involves cycles of violation and patching when work experiences fail to match self-concept. Identity work is mobilized to repair the breach when what one does diverges from who one understands oneself to be. Lamb and Davidson (2005) extended this framework to information technology, showing that IT artifacts do not simply augment professional identity; they actively reconfigure it, sometimes in ways users do not recognize. Beane (2019) provided

the most proximate precedent: when robotic surgery automated aspects of surgical skill, surgical trainees engaged in “shadow learning” (i.e., finding unofficial ways to develop competence that the formal system no longer supported) to maintain their identities as competent future surgeons.

Beyond earlier technology disruptions in 1990s and 2000s, recently, a substantial body of scholarship has already examined how professionals renegotiate autonomy and identity under algorithmic systems. Kellogg et al. (2020) characterize algorithms as a new contested terrain of control at work; Pachidi et al. (2021) show how an analytics-based “regime of knowing” reconfigured whose expertise counted in a sales organization; and Waardenburg et al. (2022) trace how predictive-policing algorithms spawned new brokering roles that reshaped professional knowledge claims. This work establishes that algorithmic systems do not merely augment work - they redistribute the authority to know. GenAI, however, introduces a distinctive form of identity threat that this literature does not fully capture.

The systems the existing studies examine typically classify, score, or recommend, and professionals experience them as an external “other” to be resisted, brokered, or worked around. GenAI instead produces the professional’s own signature output in the first person. Where earlier algorithmic systems provoked inter-occupational contests over whose judgment counts, GenAI provokes an intra-personal one: not “who owns this jurisdiction” but “was my expertise ever real?” The threat is not that an algorithm decides in the professional’s place, but that it composes in the professional’s voice.

Unlike robotic surgery, which visibly automates physical skills, GenAI automates cognitive outputs. The tool produces work that professionals could produce, but faster and often at better quality. The boundary between “what I made” and “what the tool made” becomes porous. For academics who have built an identity around being “the person who writes well” or “the person who synthesizes complex literatures,” this porosity is existentially destabilizing. Identity crisis is not primarily about being replaced, although that fear is real. It is about becoming uncertain about whether one was ever truly an expert in the first place. If a tool can do what I spent years learning to do, was my expertise real, or was it merely the slow, inefficient performance of what an algorithm now does instantly?

This crisis operates at both the individual and collective levels. At the individual level, knowledge workers experience what Pratt et al. (2006) would recognize as identity violation: the work they are doing, or that the tool is doing through them, no longer matches their self-concept as experts. They may respond by defensively asserting their unique human value (“AI can’t replicate true creativity” or “AI lacks critical thinking”), by over-identifying with the tool (“I’m an AI-powered researcher now”), or withdrawing from specific fields’ contexts where the tool threatens their identity. None of these responses addresses the underlying problem of category confusion between using and understanding AI.

At the collective level, the crisis is overwhelming. Abbott (1988) theorized that professions maintain jurisdiction through abstract knowledge systems. When tools democratize task performance, jurisdictional boundaries are renegotiated. GenAI creates a situation where the abstract knowledge base of “AI expertise” is itself shifting rapidly (the technology changes faster than any individual can master it), while the tool simultaneously democratizes the task-performance dimension. Anyone can produce AI-augmented outputs. The result is a jurisdictional chaos. Who is the “AI expert” at the university? A computer scientist who understands trans-

former architectures? An educational technologist who knows how to effectively prompt for pedagogical scenarios? The administrator who has read several articles about AI in higher education and now speaks fluently about “the opportunities and challenges”? Each has a claim, but the claims are incommensurable in nature. They rest on different forms of knowledge, modes of demonstration, and sources of legitimacy.

Not disputing that AI literacy cannot create new forms of AI-related expertise, Barley (1996) showed that new technologies can create new expert roles that reconfigure existing hierarchies. However, Barley (1996) studied the moment when the technology stabilized sufficiently for new roles to crystallize. GenAI has not stabilized, nor is it likely to do so in the foreseeable future. The jurisdictional negotiations that Abbott described are occurring in real time, but without the fixed reference points that previously allowed professions to reach settlements. The category of “AI expert” is being constructed, contested, and potentially dissolved faster than the social processes that ordinarily regulate professional identity can operate.

#### 4. The Central Paradox

The emerging empirical literature on GenAI’s effects contains a paradox that has not been adequately reconciled, and this paradox is directly relevant to the professional identity crisis I am describing.

Brynjolfsson et al. (2023) studied 5,179 customer support agents using a GPT-based assistant and found that novice and low-skill workers improved productivity by 34 percent, while high-skill workers saw minimal gains. The AI effectively transferred best practice patterns from experts to novices through the tool itself. Based on this evidence, GenAI appears to compress the expertise gap. Tool access functionally substitutes for experience in structured, task-specific fields. If we infer from this finding, academic administrators who believe that their institution no longer needs AI consultants because everyone has ChatGPT have empirical support for that position. The tool does, in measurable ways, level the playing field between those with formal AI training and those without.

Dell’Acqua et al. (2023) provided a counterpoint. Studying 758 consultants using GPT-4 on realistic tasks, they found that for tasks within GPT-4’s capability frontier, consultants improved by 40 percent. However, on tasks outside the frontier, those using AI performed 23 percent worse than those who did not. Critically, participants could not reliably distinguish frontier-inside tasks from frontier-outside tasks. Tool access without judgment was net-negative for boundary tasks.

These findings are not contradictory in the sense that one must be wrong. They are contradictory in a deeper sense: they reveal that GenAI simultaneously compresses the measurable expertise gap and creates conditions under which users cannot recognize when their expertise is failing. This is the “jagged technological frontier” (Dell’Acqua et al., 2023) mapped onto the Dunning-Kruger effect. The tool makes novices better at tasks within its capabilities while leaving them unable to perceive the boundaries of those capabilities. The confidence derived from improved performance on frontier-inside tasks bleeds into frontier-outside tasks, where it is precisely what produces failures.

Consider the implications for academics who have begun to identify as “AI expert” based on their GenAI proficiency. On frontier-inside tasks (e.g., summarizing articles, generating discussion questions, and drafting emails), their performance improves

markedly. They receive positive reinforcement from colleagues who are impressed by their fluency with new technology. Their provisional identity as “someone who knows AI” is validated. When asked to make consequential decisions about AI (whether to adopt a particular tool that raises privacy concerns, how to evaluate a vendor’s claims about their AI tools’ capabilities, whether a student’s use of GenAI constitutes plagiarism or legitimate assistance), they are operating on frontier-outside terrain without recognizing it. And they cannot tell. The metacognitive skills that would signal “I don’t know enough to make this judgment” have been eroded by the tool’s fluency on easier tasks.

This is not merely an individual issue. It scales. Organizations that hire and promote based on Gen AI-augmented output rather than demonstrated understanding of AI are building decision-making structures on a foundation of simulated expertise. When the jagged frontier is crossed, and it will be, the organization will have no internal capacity to recognize or correct the failure. An academic administrator who believes that their institution’s AI needs are solved will not perceive the moment when they are not. They lack the conceptual understanding necessary to perceive it. The tool that made them feel competent also made them incapable of recognizing their incompetence.

## 5. Implications for Higher Education and Professional Practice

The confusion regarding the distinction between GenAI tool use and AI expertise has immediate and pressing consequences for higher education, both as a site of knowledge production and as an institution that certifies expertise.

Consider first the implications for educators. Krammer (2023) argued that AI’s entry into management education demands critical reflection on what it means for how and what we teach. The same is true for other disciplines. Faculty members who lack a conceptual understanding of AI are nonetheless being asked to make consequential pedagogical decisions about it: should students be permitted to use GenAI on assignments? If so, what are the attribution requirements? How should curricula be redesigned to account for AI’s capabilities? Making these decisions well does not require an instructor to build language models or master transformer mathematics; that is developer-level expertise, and demanding it of every humanities or social-science instructor would be neither realistic nor necessary.

What it does require is critical conceptual literacy that comes from genuine curiosity: a working grasp of what these systems are, how they generate text, why they hallucinate, where their training data come from, what biases and limits follow, and what the pedagogical and ethical stakes are. Therefore, an educator with only functional facility with GenAI - who can prompt effectively but cannot reason about its limitations, failure modes, or appropriate uses is not yet equipped to make these decisions, even though the fluency of the tool can make them feel equipped.

This problem is compounded by institutional incentives. Universities are under pressure to demonstrate that they are “embracing AI” and preparing students for an AI-augmented workforce. This pressure creates a demand for visible AI initiatives, such as workshops on prompt engineering, AI literacy requirements in curricula, and statements of AI strategy. Yet, the content of these initiatives often reflects the same category confusion I have described. Prompt engineering workshops teach functional use without conceptual understanding of AI itself. AI literacy requirements are designed by committees whose members learned about

AI from the same tools they are now teaching students to use. The blind lead the blind, and both believe that they can see.

On a constructive criticism route, Selwyn (2022) warned that AI in education is surrounded by hype that obscures its actual pedagogical limitations, and that educators should resist technosolutionism. The warning is apt but insufficient. The problem is not merely that educators may be overly optimistic about AI’s capabilities. The problem is that they cannot reliably distinguish between warranted and unwarranted optimism because they lack the conceptual framework to make the distinction. They are operating on the jagged frontier without knowing where the frontier lies.

Mollick and Mollick (2023) offered a more optimistic view, arguing that AI can serve as an effective pedagogical partner when thoughtfully integrated. I do not dispute this. But the key phrase is “thoughtfully integrated.” Thoughtful integration requires two forms of instructional expertise. The first is knowing what students need to learn, how they learn it, and how a tool can support rather than undermine that learning. Second is AI expertise which includes knowing what the tool actually does, what it cannot do, and when its use is inappropriate. The educator who possesses the first form of expertise but not the second is not equipped for thoughtful integration. They are equipped for unreflective adoption, which is precisely what Selwyn (2022) warns against.

The implications extend beyond pedagogy to the institutional certification of expertise itself. Higher education is in the business of determining who knows what and certifying that knowledge through degrees and credentials. Abbott’s (1988) jurisdictional theory reminds us that this certification function is not merely administrative; it is the mechanism through which professions maintain their social contract. When a university certifies that someone has “AI expertise”, through a degree program, a certificate, or simply by hiring them into a role that implies such expertise, it is making a jurisdictional claim. It is telling the world that this person can be trusted to make AI-related judgments.

What happens when the basis for that certification is functional use rather than conceptual understanding? We are already seeing the emergence of “AI certificates” and “AI credentials” that require little more than demonstrated proficiency with GenAI tools. These credentials are not fraudulent in any straightforward sense. They genuinely certify that the holder can use the tools effectively. Yet they are misleading because they imply a broader competence that the holder does not possess. The market for AI expertise is being flooded with credentials that signal simulated contributory expertise rather than the genuine AI expertise. Because those issuing the credentials often share the same category confusion as those receiving them, there is no gatekeeping function to maintain the distinction.

This is not an argument against GenAI in education. The tools are genuinely useful, and their productivity benefits are real and well documented. The argument is against the category error that treats tool proficiency as equivalent to specific field expertise. This error is not merely semantic. It has consequences for who gets hired, whose judgment is trusted in faculty meetings, and how institutions allocate resources for AI initiatives. It also has consequences for students who may graduate with credentials that overstate their actual capabilities and find themselves unprepared for the frontier-outside tasks they will encounter in professional practice.

### 5.1. What Can Be Done?

I offer two modest proposals, not as solutions but as starting points for a more honest conversation.

First, institutions should explicitly distinguish between "AI literacy" and "AI expertise" in their formal communications, job descriptions, and credentialing practices. This distinction should be operationalized: AI literacy means demonstrated ability to use AI tools effectively for specified tasks. AI expertise means demonstrated understanding of AI fundamentals, including model architectures, training processes, limitations, failure modes, and ethical implications. The two are not mutually exclusive, one can possess both, but they are not equivalent, and possession of the first should not be mistaken for possession of the second. To make this distinction usable rather than rhetorical, institutions can adopt a three-tier model that ties each level to observable evidence and to the decisions it licenses:

**Table 1.** Three-tier model from mere AI literacy to AI expertise

| Tier                                                        | What it is                                                                                                                     | Evidence                                                                                                        | Decisions it licenses                                              |
|-------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------|
| 1. Mere AI literacy (functional use – first step)           | Can prompt effectively and integrate outputs into a workflow                                                                   | Completed tasks; a portfolio of outputs                                                                         | Using tools for individual and institutions' own work              |
| 2. Critical AI literacy (governance literacy – second step) | Understands, at a conceptual level, how models generate text, why they fail, their data and bias limits, and the ethics of use | A written critique of a tool's limitations; a sample use/attribution policy; a reasoned vendor-claim assessment | Setting course policy, advising on adoption, redesigning curricula |
| 3. AI expertise (technical – third step)                    | Understands architectures and training; can build or audit systems                                                             | Technical artifacts; audits; model development                                                                  | System design and technical assurance                              |

The tiers are cumulative in credibility but not in prerequisite: a historian or any educator can reach Tier 2 without Tier 3 (Table 1). This is to say that AI-expertise require iterative acts starting from tier 1, just mere literacy to aid with functional use to tier 3 if necessary depending with the context. Operationalized, this means a search or committee filling a role such as "AI-in-teaching coordinator" should specify the tier and its evidence rather than writing "AI experience required." For a governance role, the posting would require Tier 2 evidence - say, a two-page critique of a named tool's failure modes and a draft classroom-use policy not a portfolio of prompts, which demonstrates only Tier 1.

For non-technical departments, conceptual AI training need not become a new course or add credits. A "module, not a major" approach embeds a short unit - three sessions - into an existing required methods, writing, or capstone course: one session on how language models generate text and why they hallucinate; one on evaluating reliability and provenance in the student's own discipline; one on ethics and appropriate use. The assessment is a single exercise in which students critique a GenAI-generated artifact from their field. This delivers Tier 2 literacy without expanding the credit load.

Second, institutions should cultivate epistemic humility around AI. The greater risk lies not with the person who admits they do

not understand AI. It is the one who believes they understand because they can use ChatGPT effectively. Creating organizational cultures where saying "I don't know whether this AI-generated output is reliable" is valued at least as highly as producing fluent outputs is essential. This requires leadership modeling. Academic administrators, in particular, must resist the temptation to present themselves as AI experts based on tool proficiency. They must model the distinction between use and understanding.

## 6. Conclusion

This positional paper has argued that the distinction between using AI and understanding it is where an emerging crisis of professional identity takes root. Individually, knowledge workers confront the simulacrum of their own expertise; collectively, the very category of "AI expert" dissolves into jurisdictional chaos as anyone with prompt proficiency claims the mantle. The empirical record confirms the paradox: AI compresses measurable performance gaps while simultaneously rendering users unable to perceive when their competence fails a jagged frontier with direct consequences for how higher education certifies expertise and how organizations allocate decision-making authority. Expertise has always required not only knowing but being recognized as one who knows. AI disrupts this recognition economy by democratizing the performance while leaving the underlying knowledge untouched. The path forward demands research into whether sustained use breeds confidence or genuine competence, and restraint in deploying the label "AI expert" until we have better answers. The tools are not the problem. The category error that mistakes their fluency for our own understanding is. Recognizing the distance between the appearance of expertise and its substance is the first step toward a more honest engagement with what these tools can and cannot do - and with what we can and cannot claim to know.

## REFERENCES

- [1] Abbott, A. (1988). *The system of professions: An essay on the division of expert labor*. University of Chicago Press.
- [2] Alvesson, M., & Willmott, H. (2002). Identity regulation as organizational control: Producing the appropriate individual. *Journal of Management Studies*, 39(5), 619–644. <https://doi.org/10.1111/1467-6486.00305>
- [3] Barley, S. R. (1996). Technicians in the workplace: Ethnographic evidence for bringing work into organizational studies. *Administrative Science Quarterly*, 41(3), 404–441. <https://doi.org/10.2307/2393937>
- [4] Beane, M. (2019). Shadow learning: Building robotic surgical skill when approved means fail. *Administrative Science Quarterly*, 64(1), 87–123. <https://doi.org/10.1177/0001839217751692>
- [5] Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency (FAccT '21)* (pp. 610–623). Association for Computing Machinery. <https://doi.org/10.1145/3442188.3445922>
- [6] Brynjolfsson, E., Li, D., & Raymond, L. R. (2023). *Generative AI at work* (NBER Working Paper No. 31161). National Bureau of Economic Research. <https://doi.org/10.3386/w31161>

- [7] Collins, H., & Evans, R. (2007). *Rethinking expertise*. University of Chicago Press.
- [8] Dell'Acqua, F., McFowland, E., III, Mollick, E. R., Lifshitz-Assaf, H., Kellogg, K., Rajendran, S., Kraye, L., Candelon, F., & Lakhani, K. R. (2023). *Navigating the jagged technological frontier: Field experimental evidence of the effects of AI on knowledge worker productivity and quality* (Harvard Business School Technology & Operations Management Unit Working Paper No. 24-013). <https://doi.org/10.2139/ssrn.4573321>
- [9] Dreyfus, H. L., & Dreyfus, S. E. (1986). *Mind over machine: The power of human intuition and expertise in the era of the computer*. Free Press.
- [10] Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., Baabdullah, A. M., Koohang, A., Raghavan, V., Ahuja, M., Albanna, H., Albashrawi, M. A., Al-Busaidi, A. S., Balakrishnan, J., Barlette, Y., Basu, S., Bose, I., Brooks, L., Buhalis, D., ... Wright, R. (2023). "So what if ChatGPT wrote it?" Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *International Journal of Information Management*, 71, Article 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- [11] Floridi, L., & Chiriatti, M. (2020). GPT-3: Its nature, scope, limits, and consequences. *Minds and Machines*, 30(4), 681–694. <https://doi.org/10.1007/s11023-020-09548-1>
- [12] Ibarra, H. (1999). Provisional selves: Experimenting with image and identity in professional adaptation. *Administrative Science Quarterly*, 44(4), 764–791. <https://doi.org/10.2307/2667055>
- [13] Kellogg, K. C., Valentine, M. A., & Christin, A. (2020). Algorithms at work: The new contested terrain of control. *Academy of Management Annals*, 14(1), 366–410. <https://doi.org/10.5465/annals.2018.0174>
- [14] Krammer, S. M. S. (2023). Navigating the AI revolution: The case for intelligent management education. *Management Learning*, 54(5), 757–767. <https://doi.org/10.1177/13505076231197384>
- [15] Kruger, J., & Dunning, D. (1999). Unskilled and unaware of it: How difficulties in recognizing one's own incompetence lead to inflated self-assessments. *Journal of Personality and Social Psychology*, 77(6), 1121–1134. <https://doi.org/10.1037/0022-3514.77.6.1121>
- [16] Lamb, R., & Davidson, E. (2005). Information and communication technology challenges to scientific professional identity. *The Information Society*, 21(1), 1–24. <https://doi.org/10.1080/01972240590895883>
- [17] Long, D., & Magerko, B. (2020). What is AI literacy? Competencies and design considerations. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems* (pp. 1–16). Association for Computing Machinery. <https://doi.org/10.1145/3313831.3376727>
- [18] Mollick, E. R., & Mollick, L. (2023). *Assigning AI: Seven approaches for students, with prompts* (Wharton School Working Paper). <https://doi.org/10.2139/ssrn.4475995>
- [19] Ng, D. T. K., Leung, J. K. L., Chu, S. K. W., & Qiao, M. S. (2021). Conceptualizing AI literacy: An exploratory review. *Computers and Education: Artificial Intelligence*, 2, Article 100041. <https://doi.org/10.1016/j.caeai.2021.100041>
- [20] Pachidi, S., Berends, H., Faraj, S., & Huysman, M. (2021). Make way for the algorithms: Symbolic actions and change in a regime of knowing. *Organization Science*, 32(1), 18–41. <https://doi.org/10.1287/orsc.2020.1377>
- [21] Pratt, M. G., Rockmann, K. W., & Kaufmann, J. B. (2006). Constructing professional identity: The role of work and identity learning cycles in the customization of identity among medical residents. *Academy of Management Journal*, 49(2), 235–262. <https://doi.org/10.5465/amj.2006.20786060>
- [22] Risko, E. F., & Gilbert, S. J. (2016). Cognitive offloading. *Trends in Cognitive Sciences*, 20(9), 676–688. <https://doi.org/10.1016/j.tics.2016.07.002>
- [23] Selwyn, N. (2022). The future of AI and education: Some cautionary notes. *European Journal of Education*, 57(4), 620–631. <https://doi.org/10.1111/ejed.12532>
- [24] Svenningsson, S., & Alvesson, M. (2003). Managing managerial identities: Organizational fragmentation, discourse and identity struggle. *Human Relations*, 56(10), 1163–1193. <https://doi.org/10.1177/00187267035610001>
- [25] Waardenburg, L., Huysman, M., & Sergeeva, A. V. (2022). In the land of the blind, the one-eyed man is king: Knowledge brokerage in the age of learning algorithms. *Organization Science*, 33(1), 59–82. <https://doi.org/10.1287/orsc.2021.1544>