



After the Policy

What Faculty Actually Need in the Age of AI: Teaching, Assessment, and Scholarship

A Kolo AI Solutions White Paper

Dr. Samuel Burbanks

August 2026

This paper is written for faculty and is distributed freely. Share it with colleagues in unmodified form.

1. The Situation You Are In

If you teach at a college or university in 2026, your institution has probably given you something about AI. A policy. A syllabus template. A list of what students may and may not do. Perhaps a guidance document with suggestions that, read closely, imply you should redesign your courses on your own time. What your institution has probably not given you is help: time, training that fits your discipline, working examples, or a colleague down the hall who has figured any of this out. You received governance. You did not receive capability. If that describes your situation, this paper was written for you.

The research documents this imbalance with some precision. An analysis of 116 research universities found that 63 percent now encourage generative AI use and more than half provide sample syllabi, yet the same study warns that the guidance itself has become burdensome, because it suggests substantial revisions to teaching practice without supplying the support to make them [1]. A synthesis of 125 empirical studies names the result the curriculum integration paradox: institutional investment in AI-related faculty development correlates almost not at all with actual pedagogical change, while roughly 73 percent of students have adopted the tools and faculty resistance persists [2]. A systematic review of educators' trust in generative AI locates the largest gap not in training hours but in missing institutional leadership support [3]. In plain terms: students moved, administrations wrote rules, and faculty were left standing between the two with a policy in one hand and an unchanged workload in the other.

Beneath the policy problem sit three quieter gaps that conversations with faculty surface again and again. There is a usage gap: many professors have barely used these tools at all, or use them at the most basic level, which means every policy discussion happens in the abstract, about a technology most of the room has not handled. There is a teaching gap: the question of whether AI could serve instruction, rather than merely threaten it, has hardly been opened, because nobody was given time or examples. And there is a scholarship gap: most faculty have only an inkling of what the tools could do for their own research and writing, which is the arena where they would learn the technology's real character fastest and at the lowest stakes. These three gaps are connected, and they are the outline of this paper: a professor who has used the tool seriously (section 3) can redesign assessment from knowledge rather than fear (section 4), can decide where it belongs in their teaching (section 5), can put it to work in their scholarship (section 6), and can tell their institution precisely what support is owed (section 7).

One more finding belongs at the front of this paper, because it concerns how you may be feeling about all of this. The research on faculty and AI is unusually clear that resistance is not ignorance. Faculty hesitation reflects legitimate concerns about accuracy, academic integrity, and professional identity, and it responds to structured, humane support rather than exhortation [4]. If you have been made to feel that your skepticism is a deficiency, the evidence is on your side. Skepticism is the correct starting posture toward a tool that fabricates. The purpose of this paper is not to talk you out of your judgment. It is to give your judgment better material to work with: what the assessment crisis actually is, what these tools actually do, what documented options exist for redesigning assessment, how instructors are using AI to teach rather than merely to police, how it can serve your own scholarship, and what you and your colleagues are entitled to ask of your institution.

2. What the Brown Exam Actually Revealed

In spring 2026, an economics professor at Brown University gave a take-home midterm to 86 students. The average was 96 percent, in a course whose historical averages ran between 65 and 80. Suspecting widespread AI use, he announced that the midterm would be voided if final exam performance differed significantly, and gave the final in class. Eighteen students dropped the course. Nine stayed enrolled but did not sit the final. The in-class average was 48.6 percent, and nineteen students failed [5]. The story circulated widely, on NPR and elsewhere, because every professor who heard it recognized it.

The tempting reading is moral: students cheat now. A more useful reading, argued even within the higher education press, is that the episode is less about character than about measurement [6]. A take-home written product no longer demonstrates what it once demonstrated. That reading matters because it changes the response. If the problem is morality, the answers are surveillance and punishment, and the evidence says both fail. Experienced markers cannot reliably distinguish AI-assisted work from unassisted work, even when they believe they can, and suspicion itself distorts grading toward false accusations and false clearances [7]. Detection software is unreliable enough that a leading analysis concludes reliance on it is misaligned with the educational landscape [8]. Roughly a third of surveyed students have used a chatbot on an assessment, and many do not regard it as a breach at all [9]. Bans have proven ineffective in practice [10].

If the problem is measurement, the answer is design, and here the research offers real direction rather than consolation. It also offers one uncomfortable correction: the popular remedy of making assessment more authentic, closer to real-world tasks, is not by itself a shield. Controlled experiments found that generative AI completes authentic assessments well enough to pass expert scrutiny, leading the authors to conclude that the sector must shift emphasis from written outputs toward assessment that is performative, synchronous, and social [7]. The next section translates that finding, and others like it, into moves you can actually make.

3. A Working Knowledge of the Tool

You cannot design around a tool you have not used, and you cannot regulate a tool you cannot predict. The research finding here is encouraging: hands-on experience with AI's limitations, more than any briefing or policy, is what produces realistic expectations and calibrated trust [30]. An hour of deliberate experimentation, trying to make the tool fail in your own discipline, teaches more than a semester of guidance documents.

A serviceable mental model: large language models are fluent transformers of language, not databases of truth. They are strong at drafting, summarizing, restructuring, explaining a concept at a chosen level, translating between registers, generating examples and counterexamples, and applying stated criteria to a text. They have historically been weak precisely where academic work is strictest: they fabricate. In a social work classroom exercise built around this property, students found that only 38.5 percent of the citations in an AI-generated literature review were fully accurate, and the assignment of verifying them became the

lesson [29]. That figure describes the tools as most students first met them; current frontier models with source retrieval fabricate far less often, though not never, and the improvement is uneven across tiers and tasks, which is why the calibration experiments below matter more than any fixed claim about capability. The models also overclaim on mixed evidence, flatten nuance, and know nothing about your students, your syllabus, or what happened in Tuesday's seminar unless you tell them. There is also an emerging caution about what over-reliance does to the user, a phenomenon some researchers call cognitive debt, which is as relevant to us as to our students [36].

If you have not yet spent that hour, here is a structured version of it, five experiments that map the terrain of your own discipline. A caution before you run them: the results depend heavily on which tool you use, at which tier, and with which features enabled, and these systems are improving fast enough that failure modes which defined them a year ago have receded unevenly. Current frontier models with web access fabricate citations far less often than the versions that shaped most faculty's first impressions; free tiers and older models still fail in the old ways; and your students use every tier there is. The experiments are therefore not a demonstration that the tools are weak. They are a calibration of the specific tool in front of you, this month, and they are worth rerunning on both a frontier model and the free tier your students most likely use, because the gap between the two results is itself a finding about your classroom. First, give the model your best essay prompt or problem set and grade its answer honestly against your rubric; this is the exposure test from section 4, and the result will surprise you in one direction or the other. Second, ask for a literature review on a topic you know deeply, then check every citation against the actual sources. On a current frontier model with search enabled you may find few errors or none; on a free tier you will likely meet fabrication personally. Verification survives as a discipline either way, because the errors that remain concentrate exactly where scholars live, in page-specific quotations, obscure subfields, and paywalled or non-English sources, and one fabricated reference costs a professor more than a hundred clean ones earn. Third, ask it to explain a concept from your field at three levels, to a first-year student, to a senior major, and to a colleague. The explanations will mostly be strong, often impressively so; the discipline is reading every word, because the errors that remain are confident and plausible rather than obvious. Fourth, give it a piece of your own writing and ask for criticism against named standards; you will see the usefulness of a tireless first reader, and if you work with it across a long session, you will also notice that its grip on earlier material loosens late in an extended conversation in ways the opening exchanges do not predict. Fifth, map the boundary of what it knows. With web access, current tools can report accurately on your institution's public face, a boundary that is moving year by year. What does not move is the unwritten: your current students, what happened in Tuesday's seminar, the texture of your program that lives in no document. Knowing which boundary is which is the point. Two hours of this and you will hold a better working theory of the tool than most of the guidance documents you have received.

Two practical implications follow. First, the competence you need is not programming. Competency research across learner groups finds that what working professionals need is interpretation, error detection, and judgment about when to rely on AI output [31], and a twelve-competency scaffold now exists that runs from basic operation through prompting to ethical judgment [32]. Second, the failure modes are design material. Every weakness in the previous paragraph is something an assessment can be built to expose, and

every strength is something your teaching can borrow. The professor who knows exactly what the model does badly is the professor the model cannot surprise.

4. Redesigning Assessment Without Burning Down Your Course

Start with a diagnostic the world's leading universities now recommend to their own faculty: run your assessment through a frontier model and see what grade it earns [11]. If the tool scores an A, the assessment is measuring the tool. That test costs an evening and tells you which of your assessments are exposed and which are already sound. Nothing else in this section needs to happen until you know that.

From there, the literature offers a menu rather than a mandate, and the menu is worth internalizing because most items strengthen learning independently of AI. Reviews of assessment in AI-infused environments converge on process-based and staged assessment: graded drafts, annotated bibliographies, research logs, and multi-stage submissions that make the trajectory of the work visible rather than only its endpoint [12]. They converge on presence: oral defenses of written work, in-class writing components, and synchronous performance, the assessment forms the authentic-assessment experiments found most resistant to delegation [7, 12]. They converge on assessment that cultivates self-regulated learning and integrity as outcomes in their own right rather than policing them as compliance [13]. And they converge on involving students in the reform itself, since students and educators broadly agree on which assessments AI has compromised, and students accept redesign far better when they helped shape it [14].

A second family of moves brings the tool inside the assessment on your terms. The AI Assessment Scale gives you a five-level vocabulary, from no AI permitted to full AI with disclosure, that can be assigned per task rather than per course; a university-wide pilot reported a significant reduction in misconduct cases alongside improved student engagement [10]. AI-integrated assignments make the model an object of critique: students generate output and then must find its errors, test its citations, and improve it, which builds exactly the critical AI literacy the workplace will demand of them [16, 29]. Experiential designs use the tool as a support agent inside authentic tasks rather than as a ghostwriter of them [15], and design models now exist for building assignments that preserve student creativity and ownership [17]. A framework built from interviews with 61 faculty organizes the whole space into four honest postures: against, avoid, adopt, and explore [18]. There are assessments where the right answer is against, and in-class conditions protect them [19]. The skill is matching the posture to the learning outcome instead of applying one posture to everything.

A worked example makes the menu concrete. Take the most exposed assignment in the curriculum: the take-home analytical essay. The redesigned version might run as follows. Week one, students submit a proposal and annotated bibliography, with the annotations written in class in fifteen minutes, enough to show the sources have actually been read. Week three, a full draft is submitted along with a short reflection on what the student found hardest, and the draft is workshopped with peers. Week five, the final essay is submitted with a one-page process memo, and each student gives a five-minute oral defense: two questions from you, chosen after skimming their draft history, that anyone who wrote the paper could answer easily and anyone who did not could not. If you wish to bring the tool inside, add a declared-use tier: students

may use AI for brainstorming and grammar at level two of the scale, must disclose the prompts used, and know that the oral defense is where the grade lives. Notice what happened to the integrity problem: it did not need policing, because the design made delegation unprofitable. Notice also the cost: the redesign trades grading time at the end for structure throughout, which is precisely the kind of shift that requires the institutional time and support argued for in section 7.

Finally, say what you decided. The first policy responses of leading universities worldwide converged on one practice above all: clear communication of permitted AI use in the course syllabus [20]. A serviceable syllabus statement can be three sentences long: which uses of AI are permitted in this course and at what level, what must be disclosed and how, and why the policy serves the learning goals of the course rather than the convenience of enforcement. Students respect the third sentence more than the first two. Whatever mix you choose, students should be able to read it, and the act of writing it will clarify it.

5. Teaching With AI: The Abandoned Half

Governance answered one question: what students may not do. It left unanswered the question that would actually repay your attention: what the tool can do for your teaching. This is the half where leadership has largely abandoned faculty, and it is also the half with the friendliest evidence.

The instructor-facing uses are the natural on-ramp because they carry the lowest stakes: no student data, no integrity questions, full faculty control over the output. Documented practice includes drafting and differentiating course materials, generating worked examples and case variations, producing custom visuals for complex processes, building rubrics and question banks, and organizing course logistics, with discipline-specific playbooks now appearing that include example prompts faculty can lift directly [25, 26]. Reviews of real implementations consistently find teaching material generation, evaluation support, and feedback among the applications that actually deliver [26]. None of this changes what happens between you and your students. It changes how much of your Sunday is spent producing the third version of a handout.

The pattern behind effective instructor prompting is worth stating once, because it transfers across every use above. Give the model your context, your constraints, and your standards, then ask for a draft you will edit. Not 'write a quiz on cellular respiration' but 'here are my lecture notes and learning objectives for this unit; write eight multiple choice questions at two difficulty levels, each with a plausible-wrong distractor based on a misconception you can name.' Not 'make a rubric' but 'here is the assignment and here are three anonymized past submissions I graded A, B, and C; draft a rubric that would reproduce those judgments and tell me where my grading looks inconsistent.' The second forms work because they put your expertise inside the prompt, which is exactly what a teaching assistant would need from you, and it is the honest way to think about the tool: an infinitely patient, occasionally wrong assistant who has never met your students. There are also student-facing uses worth knowing about, tutoring dialogues, role-played historical figures and case characters, and real-time feedback on technique, documented in the discipline playbooks [25], but they belong to a second stage, after your own fluency is established and with the design-intent guardrails below.

Feedback is where the evidence becomes striking. In a physics education study, instructors judged roughly 70 percent of AI-drafted feedback on student conceptual responses usable with minor or no modification, and students rated the AI-drafted feedback as more useful than the human-written version while barely being able to tell which was which [22]. A controlled six-week study of language learners found no difference in learning outcomes between students receiving AI-generated and tutor-written feedback, with the authors recommending a blended approach [23]. A systematic review concludes that automated feedback systems can reduce routine grading load precisely so that instructors can spend their attention on the complex, relational parts of teaching [21]. The pattern across these studies is not that AI replaces your feedback. It is that AI produces a competent first draft of the routine feedback, and you remain the editor, the standard, and the human who knows the student.

Two guardrails keep this honest. The first is design intent: frameworks for human-centered use insist the tool scaffold thinking rather than substitute for it, in your students' work and in yours, and courses that integrated AI without off-loading learning did it by requiring students to think first and reflect after [27, 28]. The second is that capability grows with structured practice, not exposure: workshop-based programs measurably increased faculty confidence, in one series by around 30 percent, and longer programs improved instructors' actual ability to design AI-integrated learning [24, 28]. Which returns us to the institutional question, because a workshop series is exactly the kind of thing faculty should not have to improvise alone.

6. Your Own Scholarship

The quiet half of the faculty predicament is that most professors have only an inkling of what these tools can do for their own research and writing. Studies of professionals who use generative AI regularly find something more interesting than laziness: they do not reduce effort, they reallocate it, spending less on production and more on steering, verifying, and judging [34]. That is a fair description of what scholarly AI use looks like when it works.

The productive uses map onto the parts of scholarship that are labor rather than thought: generating candidate structures for a paper and arguing against your own framing, tightening prose you have already written, drafting the routine sections of grant applications and IRB narratives, writing and debugging analysis code, and triaging literature by summarizing papers you supply to it. The boundaries are equally concrete. Never cite a source you have not opened; the fabrication rate documented earlier applies to your bibliography as much as your students' [29]. Verify every factual claim that survives into your text. Disclose use according to your journal's and funder's policies, which now exist almost everywhere and vary meaningfully. Keep unpublished data, identifiable human subjects material, and anything covered by FERPA or HIPAA out of consumer tools whose data handling you have not confirmed. Authorship, and the responsibility that attaches to it, remains entirely yours.

A concrete workflow shows the shape of disciplined use. Suppose you are returning to a literature you have not touched in three years. Gather the fifteen most recent papers yourself, from the databases you already trust, and supply the PDFs or abstracts to the model rather than asking it to find sources. Ask for a structured summary of each against questions you specify: sample, method, central finding, stated limitation. Read its

summaries against the three papers you know best; where it flattens or overstates, you have calibrated it. Then ask it to map disagreements across the set and propose the three questions the literature has not answered. Every claim that survives into your manuscript gets checked against the source it came from, and the sources it proposed that you did not supply get checked for existence before anything else. The model did the clerical labor; the scholarship, the judgment about what matters and what is true, never left your hands. The same division of labor applies to grant boilerplate, methods sections, and analysis code, each with the same rule: it drafts, you verify, you own.

There is a second return on this competence: credibility. The professor who uses these tools seriously in their own scholarship walks into the classroom knowing what the tools are, which is the precondition for every design decision in sections 4 and 5, and students extend more trust on this subject to faculty who plainly know the terrain.

7. What to Ask of Your Institution

Faculty frustration at being handed governance without support is justified, but frustration is not a strategy. The research gives faculty something better: a specific, evidence-backed list of what institutions owe the people they expect to carry this change. Faculty senates, unions, and departmental leadership can carry these asks upward as documented findings rather than complaints.

The evidence supports asking for the following. Faculty development treated as infrastructure with dedicated time, not optional enrichment appended to full workloads [4]. Training that is hands-on and discipline-specific, since competencies differ by role and generic sessions demonstrably miss [31, 4]. Psychologically safe space to experiment and fail, because fear of looking incompetent is a documented adoption barrier that punitive cultures make worse [4]. Named peer champions and communities of practice, the mechanism by which readiness actually spreads through an organization [30]. Visible leadership engagement, since its absence is the single largest trust gap the literature identifies [3]. Institutional resourcing of assessment redesign, because reviews are unanimous that assessment reform is an institutional project requiring policy, values, and support systems, not a task to be delegated silently to individual syllabi [13, 33]. And honest acknowledgment that guidance without time is not support but burden shifting [1].

Notice what this list is not. It is not a demand that skeptical colleagues be converted, and it is not a request for more documents. It is a request that the institution invest in its faculty at the level it has already invested in its rules.

While the institutional ask moves at institutional speed, departments do not have to wait. The mechanisms the research identifies as effective, peer champions, communities of practice, and psychologically safe experimentation [30, 4], are all available at the departmental level for the cost of a standing hour. Three colleagues who run the section 3 experiments and compare notes over lunch constitute a community of practice. One colleague who redesigns one assignment and reports honestly on what worked and what did not is a champion. A department chair who says out loud that experimenting and failing with these tools is legitimate professional activity has supplied more psychological safety than most provost communications to date. The evidence suggests these small structures are not a consolation prize while waiting for

institutional support; they are the mechanism through which readiness actually spreads, and departments that build them will be the ones ready to use institutional support when it finally arrives.

8. The Whole Person at the Front of the Room

A closing word about what all of this is for. Teaching is relational work. The research on AI and meaningful work is direct on this point: the same technology can enhance or diminish the meaning people find in their work depending entirely on how it is introduced and who controls it [35]. Deployed carelessly, from above, as rules plus silence, AI makes teachers feel managed. Deployed by the teacher, as a tool under their judgment that returns hours to the parts of the work only a human can do, the same technology serves the relationship between teacher and student instead of intruding on it. The difference is not in the model. It is in who holds it, and whether the institution treated its faculty as whole people or as a compliance risk.

That conviction is the foundation of the firm publishing this paper. Kolo AI Solutions is an AI consulting practice founded by career educators and practitioners, grounded in African centered organizational thought, serving education, social services, nonprofit, and government institutions. We wrote this paper for free distribution because the research says the first thing faculty need is not a vendor but an honest map. If your department, program, or institution wants to go further, we offer workshops, faculty development designed with faculty rather than at them, and structured assessments of where your institution actually stands with AI. Reach us at koloaisolutions.com or samuel@koloaisolutions.com.

Dr. Samuel Burbanks holds a Ph.D. in Education and an M.Ed. in Secondary Education. He taught for Cincinnati Public Schools for fifteen years and serves as an adjunct instructor in the College of Education at the University of Cincinnati and at Northern Kentucky University, and with Upward Bound at the University of Cincinnati. He leads the education vertical and serves as AI technical lead at Kolo AI Solutions.

References

- [1] [Generative artificial intelligence in higher education: Evidence from an analysis of institutional policies and guidelines](#) (McDonald et al., 2025, Computers in Human Behavior: Artificial Humans)
- [2] [Generative artificial intelligence in higher education: A systematic review of perceptions, implementation and pedagogical transformation](#) (Segura Altamirano et al., 2026, Review of Education)
- [3] [Exploring trust in generative AI for higher education institutions: a systematic literature review focused on educators](#) (Lelescu et al., 2025, Humanities and Social Sciences Communications)
- [4] [From resistance to readiness: faculty development as the key to AI literacy in public health](#) (Acosta, 2026, Frontiers in Public Health)
- [5] [Brown Professor Suspects Most of His Class Used AI to Cheat](#) (Inside Higher Ed, July 2026)
- [6] [That Chart From Brown Isn't Really About Cheating](#) (Inside Higher Ed opinion, July 2026)
- [7] [The impact of generative AI on academic integrity of authentic assessments within a higher education context](#) (Kofinas et al., 2025, British Journal of Educational Technology)
- [8] [Generative AI detection in higher education assessments](#) (Ardito, 2024, New Directions for Teaching and Learning)
- [9] [The rapid rise of generative AI and its implications for academic integrity: Students' perceptions and use of chatbots for assistance with assessments](#) (Gruenhagen et al., 2024, Computers and Education: Artificial Intelligence)
- [10] [The AI Assessment Scale \(AIAS\) in action: A pilot implementation of GenAI supported assessment](#) (Furze et al., 2024, ArXiv)
- [11] [Generative AI tools and assessment: Guidelines of the world's top-ranking universities](#) (Moorhouse et al., 2023, Computers and Education Open)
- [12] [Redefining student assessment in AI-infused learning environments: a systematic review of challenges and strategies for academic integrity](#) (Ncube et al., 2025, AI and Ethics)
- [13] [A scoping review on how generative artificial intelligence transforms assessment in higher education](#) (Xia et al., 2024, International Journal of Educational Technology in Higher Education)
- [14] [Educator and Student Perspectives on the Impact of Generative AI on Assessments in Higher Education](#) (Smolansky et al., 2023, ACM Learning @ Scale)
- [15] [Using Generative Artificial Intelligence Tools to Explain and Enhance Experiential Learning for Authentic Assessment](#) (Salinas-Navarro et al., 2024, Education Sciences)
- [16] [Participatory Co-Design and Evaluation of a Novel Approach to Generative AI-Integrated Coursework Assessment in Higher Education](#) (Martin et al., 2025, Behavioral Sciences)
- [17] [Redefining assessment tasks to promote students' creativity and integrity in the age of generative artificial intelligence](#) (Peters et al., 2025, International Journal for Educational Integrity)
- [18] [Redesigning Assessments for AI-Enhanced Learning: A Framework for Educators in the Generative AI Era](#) (Khlaif et al., 2025, Education Sciences)
- [19] [Ensuring academic integrity in the age of ChatGPT: Rethinking exam design, assessment strategies, and ethical AI policies in higher education](#) (Evangelista, 2025, Contemporary Educational Technology)

- [20] [AI and ethics: Investigating the first policy responses of higher education institutions to the challenge of generative AI](#) (Dabis et al., 2024, Humanities and Social Sciences Communications)
- [21] [Harnessing Generative AI \(GenAI\) for Automated Feedback in Higher Education: A Systematic Review](#) (Lee et al., 2024, Online Learning)
- [22] [Exploring generative AI assisted feedback writing for students' written responses to a physics conceptual question with prompt engineering and few-shot learning](#) (Wan et al., 2023, Physical Review Physics Education Research)
- [23] [AI-generated feedback on writing: insights into efficacy and ENL student preference](#) (Escalante et al., 2023, International Journal of Educational Technology in Higher Education)
- [24] [Integration of Generative Artificial Intelligence in Higher Education: Best Practices](#) (Cordero et al., 2024, Education Sciences)
- [25] [Teaching microbiology with generative AI: a survey-style playbook to empower microbiology faculty to integrate AI into courses](#) (Mosovsky et al., 2026, Journal of Microbiology and Biology Education)
- [26] [Generative AI Solutions for Faculty and Students: A Review of Literature and Roadmap for Future Research](#) (Marchena Sekli et al., 2024, Journal of Information Technology Education: Research)
- [27] [Incorporating generative AI into a writing-intensive undergraduate course without off-loading learning](#) (Tate et al., 2025, Discover Computing)
- [28] [A Human-Centered Learning and Teaching Framework Using Generative Artificial Intelligence for Self-Regulated Learning Development](#) (Kong et al., 2024, IEEE Transactions on Learning Technologies)
- [29] [Innovating social work pedagogy: generative AI for critical competency development in evidence-based practice](#) (Walker et al., 2026, Social Work Education)
- [30] [Making Sense of AI Limitations: How Individual Perceptions Shape Organizational Readiness for AI Adoption](#) (Übellacker, 2025, ArXiv preprint)
- [31] [A Competency Framework for AI Literacy: Variations by Different Learner Groups and an Implied Learning Pathway](#) (Chee et al., 2024, British Journal of Educational Technology)
- [32] [Generative AI Literacy: Twelve Defining Competencies](#) (Annapureddy et al., 2024, Digital Government: Research and Practice)
- [33] [Designing assessments in the generative AI era: A tailored assessment framework for ICT tertiary education](#) (Ahangama, 2026, International Journal of Educational Technology in Higher Education)
- [34] [From Effort Reduction to Effort Management: An Expectancy Theory Perspective on Professionals' Work Practices with Generative AI](#) (Memmert et al., 2025, Business & Information Systems Engineering)
- [35] [The Ethical Implications of Artificial Intelligence \(AI\) For Meaningful Work](#) (Bankins et al., 2023, Journal of Business Ethics)
- [36] [Artificial Intelligence in Higher Education: A State-of-the-Art Overview of Pedagogical Integrity, Artificial Intelligence Literacy, and Policy Integration](#) (Adamakis et al., 2025, Encyclopedia)