

VOL. XI ISSUE I, WINTER (March-2026)

GRR

GLOBAL REGIONAL REVIEW

GLOBAL REGIONAL REVIEW (GRR)

Title: Beyond Tool Use: Accountable Scholarly Augmentation, Responsible AI Literacy, and Publication Readiness Among Emerging Scholars

Abstract

Generative artificial intelligence (GenAI) is rapidly transforming postgraduate research communication, creating opportunities for academic assistance while also posing risks to source authenticity and intellectual stewardship. Based on an empirical multi-site meta-synthesis of N=4,280 doctoral and early-career researchers across eight major universities, this study establishes and validates the Accountable Scholarly Augmentation framework. Findings show GenAI effectively supports low-level tasks like outlining and stylistic framing. However, unguided use creates an "Overconfidence Gap," where high prose fluency masks factual errors and citation hallucinations (up to 38.2% unverified claims). Multivariate analysis reveals that punitive policy warnings fail to reduce manuscript defects. Conversely, integrated responsible AI programs significantly decrease pre-submission desk rejections from 41.8% to 12.6% ($\beta=.52, p<.001$). The study emphasizes shifting research integrity governance and graduate education from punitive policy enforcement to a developmental publication-readiness framework.

Keywords: Accountable Scholarly Augmentation, Generative AI in Higher Education, Overconfidence Gap, Academic Research Integrity, Postgraduate Writing Governance

Authors:

Sana Hussan: (Corresponding Author)

PhD Scholar, Business Administration, The University of the Potomac, Virginia, United States (US)
(Email: sanahusan143@gmail.com)

Pages: 1-14

DOI:10.31703/grr.2026(XI-I).01

DOI link: [https://dx.doi.org/10.31703/grr.2026\(XI-I\).01](https://dx.doi.org/10.31703/grr.2026(XI-I).01)

Article link: <https://grrjournal.com/article/beyond-tool-use-accountable-scholarly-augmentation-responsible-ai-literacy-and-publication-readiness-among-emerging-scholars>

Full-text Link: <https://grrjournal.com/article/beyond-tool-use-accountable-scholarly-augmentation-responsible-ai-literacy-and-publication-readiness-among-emerging-scholars>

Pdf link: <https://www.grrjournal.com/jadmin/Auther/31rvIolA2.pdf>

Global Economics Review

p-ISSN: [2616-955X](https://doi.org/10.31703/grr.2026(XI-I).01) e-ISSN: [2663-7030](https://doi.org/10.31703/grr.2026(XI-I).01)

DOI(journal):10.31703/grr

Volume: XI (2026)

DOI (volume):10.31703/grr.2026(XI-I)

Issue: I Winter (March-2026)

DOI(Issue):10.31703/grr.2026(XI-I)

Home Page

www.grrjournal.com

Volume: XI (2026)

<https://www.grrjournal.com/Current-issue>

Issue: I-Winter (March 2026)

<https://www.grrjournal.com/issue/11/1/2026>

Scope

<https://www.grrjournal.com/about-us/scope>

Submission

<https://humaglobe.com/index.php/grr/submissions>

Scan the QR to visit us



Google
scholar



Citing this Article

Article Serial	01
Article Title	Beyond Tool Use: Accountable Scholarly Augmentation, Responsible AI Literacy, and Publication Readiness Among Emerging Scholars
Authors	Sana Hussan
DOI	10.31703/grr.2026(XI-I).01
Pages	1-14
Year	2026
Volume	XI
Issue	I

Referencing & Citing Styles

APA	Hussan, S. (2026). Beyond Tool Use: Accountable Scholarly Augmentation, Responsible AI Literacy, and Publication Readiness Among Emerging Scholars. <i>Global Regional Review</i> , 11(1), 1-14. https://doi.org/10.31703/grr.2026(XI-I).01
CHICAGO	Hussan, Sana. 2026. "Beyond Tool Use: Accountable Scholarly Augmentation, Responsible AI Literacy, and Publication Readiness Among Emerging Scholars." <i>Global Regional Review</i> 11 (1):1-14. doi: 10.31703/grr.2026(XI-I).01.
HARVARD	HUSSAN, S. 2026. Beyond Tool Use: Accountable Scholarly Augmentation, Responsible AI Literacy, and Publication Readiness Among Emerging Scholars. <i>Global Regional Review</i> , 11, 1-14.
MHRA	Hussan, Sana. 2026. 'Beyond Tool Use: Accountable Scholarly Augmentation, Responsible AI Literacy, and Publication Readiness Among Emerging Scholars', <i>Global Regional Review</i> , 11: 1-14.
MLA	Hussan, Sana. "Beyond Tool Use: Accountable Scholarly Augmentation, Responsible Ai Literacy, and Publication Readiness among Emerging Scholars." <i>Global Regional Review</i> 11.1 (2026): 1-14. Print.
OXFORD	Hussan, Sana (2026), 'Beyond Tool Use: Accountable Scholarly Augmentation, Responsible AI Literacy, and Publication Readiness Among Emerging Scholars', <i>Global Regional Review</i> , 11 (1), 1-14.
TURABIAN	Hussan, Sana. "Beyond Tool Use: Accountable Scholarly Augmentation, Responsible Ai Literacy, and Publication Readiness among Emerging Scholars." <i>Global Regional Review</i> 11, no. 1 (2026): 1-14. https://dx.doi.org/10.31703/grr.2026(XI-I).01 .

Beyond Tool Use: Accountable Scholarly Augmentation, Responsible AI Literacy, and Publication Readiness Among Emerging Scholars



Sana Hussan (Corresponding Author)¹

¹ PhD Scholar, Business Administration, The University of the Potomac, Virginia, United States (US)
(Email: sanahusan143@gmail.com)

Abstract

Generative artificial intelligence (GenAI) is rapidly transforming postgraduate research communication, creating opportunities for academic assistance while also posing risks to source authenticity and intellectual stewardship. Based on an empirical multi-site meta-synthesis of N=4,280 doctoral and early-career researchers across eight major universities, this study establishes and validates the Accountable Scholarly Augmentation framework. Findings show GenAI effectively supports low-level tasks like outlining and stylistic framing. However, unguided use creates an "Overconfidence Gap," where high prose fluency masks factual errors and citation hallucinations (up to 38.2% unverified claims). Multivariate analysis reveals that punitive policy warnings fail to reduce manuscript defects. Conversely, integrated responsible AI programs significantly decrease pre-submission desk rejections from 41.8% to 12.6% ($\beta=.52, p<.001$). The study emphasizes shifting research integrity governance and graduate education from punitive policy enforcement to a developmental publication-readiness framework.

Keywords: *Accountable Scholarly Augmentation, Generative AI in Higher Education, Overconfidence Gap, Academic Research Integrity, Postgraduate Writing Governance*

Introduction

On the academic research side, generative artificial intelligence (GenAI) has infiltrated at a rate unprecedented to other shifts in education technology. Large language models (LLMs) and conversational AI assistants have become common tools for researchers, doctoral candidates, Master's students, and postdoctoral students across many aspects of research communication. Researchers use these tools to brainstorm research questions, organize and conduct literature searches, improve academic writing, suggest target journals, summarize dense theoretical text, and simulate peer feedback. This adoption is not merely a 'novelty technology', but a pragmatic and rational response to deep structural characteristics of modern higher education. Upcoming researchers face considerable institutional pressure to publish in high-impact journals, tight deadlines, limited supervision, and a host of inequalities between monolingual and multilingual scholars when publishing in a second or third language.

But the high level of "accessibility" and "fluency" in language generation with generative AI carries huge developmental and ethical challenges. Outputs from such large algorithms often feature well-written, convincing text with hidden subtractions of facts, fake reference sources, shallow theoretical thinking, and illogicalities. If early-career scholars become accustomed to using generative tools before they mature enough to develop and monitor their disciplinary ideas, they may delude themselves into thinking that mastery is even superficial mastery of their texts. If the early-career scholar considers surface competency with text as some level of intellectual control, he/she may be deluding him- or herself. The developmental



question is thus the key challenge to postgraduate research training, rather than a simply disciplinary, punitive one. Warnings about institutional policy and usage of generative AI, or outright prohibitions on such tools, don't communicate the reality: that generative AI is already a part of the academic life cycle – and that a ban on its use will simply drive usage underground, into someone's own, untracked research or sandboxing.

It makes sense that generative AI sounds appealing to newly minted researchers. Postgraduate research is highly ambiguous and cognitively demanding. Researchers need to know specific jargon related to their field, to (often clumsily!) come up with new ideas that bring together hundreds of empirical studies, to pinpoint subtle theoretical weaknesses, to formulate defensible research techniques, and to write very formally. Generative tools are emerging as an on-demand editorial service and feedback for first-generation academics or researchers who do not have many connections. In this context, GenAI serves as an informal 'hidden curriculum decoder' that can guide researchers in structuring abstracts to meet journal expectations, smooth rough edges, and even translate conversational logic into academic discourse. Those who would deny the pragmatic possibilities of these uses are not facing reality when it comes to equity issues for new scholars.

But many of these attributes of generative tools are also deceiving development pitfalls. LLMs do not provide old-style writing tools like spell checkers or reference managers by returning "facts" verified by the document; rather, given the provided context, they return plausible sequences of words stochastically. They produce text with an authoritative tone despite the lack of empirical validity and logical coherence in what they assert. An LLM producing a literature synthesis might potentially make superficial attempts to reconcile real differences in theoretical perspectives, "hallucinate" references that sound convincing or plausible, or create subtle misinterpretations of original empirical work on the part of the emerging scholar. Without this kind of tacit knowledge, often possessed by the scholar alone, the impersonated content could be taken as truth, and the scholar falls into essentially a "let the computer do the critical thinking" trap.

Responses to this phenomenon have been quite polarized at the institutional level. Others at the front lines of higher education are adopting all-encompassing policy statements with sanctions for any wrongdoing with AI. At the other end of the spectrum, techno-optimists argue that GenAI is a transformative tool that could increase productivity and streamline the process of turning literature into manuscripts. Neither position offers a pedagogical argument. Blanket prohibitions are difficult to monitor and punish, tarnish the reputations of honorable students, and are especially hard for a non-native English writer who uses language editing tools. On the other hand, failing to reflect on the documentable decline in rigor, integrity, and ideas that takes place when AI generation takes center stage is uncritical celebration. What is needed is a middle way grounded in careful, thorough empirical study and pedagogical scaffolding.

The structural changes brought about by generative AI also affect graduate faculty and dissertation supervisors. Academic advisors also have to cope with an extremely high workload and, often, have no formal training in assessing AI-assisted writing. Supervisors often find it difficult to distinguish between appropriate language polishing and inappropriate text production and, as a result, may provide conflicting feedback (and feel anxious) about postgraduate students' language. This supervisory ambiguity reinforces the need for clear institutional rules that provide concrete indicators and verification for both advisor(s) and advisee(s).

In this regard, a comprehensive framework for Responsible AI Literacy must be created – and tightly coupled with research communication and publication readiness – in universities and the research community. Responsible AI Literacy extends beyond simply being proficient in how AI works or giving "AI tips", by asking whether an emerging scholar will know when to use generated AI, when to appreciate its systematic bias or distortion, when to thoroughly review and correct works created using it, when to seek their help, and when to keep ideas entirely within his or her own control. Hence, publication-readiness should be understood as a multi-dimensional, developing ecology that includes problem formulation, identifying noted gaps, verifying sources, methodological defense, journal selection, navigating the peer review process, and public communication of research.

Operationalization in the context of an Empirical Multi-Site Meta-Synthesis: In this empirical article, the notion of Accountable Scholarly Augmentation is operationalized and evaluated for an Empirical Multi-Site Meta-Synthesis of N = 4280 postgraduates within eight main research universities. This study focuses on three main research questions concerning the factor of institutional training regimes regarding task-specific utilization, verification behavior, and manuscript defect rate:

- RQ1:** How do emerging scholars utilize generative AI across distinct research communication tasks, and where does the disconnect between perceived text fluency and actual factual/citation accuracy manifest?
- RQ2:** How do different institutional training interventions (policy warnings vs. generic tooltips vs. integrated responsible AI literacy) influence researchers' source verification habits and pre-submission desk rejection rates?
- RQ3:** To what extent does accountable scholarly augmentation mediate the relationship between responsible AI literacy and holistic publication readiness?

Literature Review

AI Literacy and Post-Graduate Scholarly Development

The idea of “AI literacy” has quickly evolved from the basic notion of “digital literacy.” According to Long and Magerko (2020), several key skills and abilities can be considered as AI literacy: critical evaluation of technology, communication skills with intelligent systems, and collaboration skills with AI in a variety of settings. Taking this further, Ng et al. (2021) envisaged AI literacy in four aspects: cognitive (understanding how AI works), affective (attitude and trust), socio-organizational (impacts and implications for society and institutions), and ethical (fairness, accountability, and privacy). In higher education, they focus on helping students critically evaluate these frameworks and their ethical implications.

AI literacy needs to be tailored especially to the requirements of academic research and publication for rising scholars. Academic research communication is not part of ordinary administrative duties; it includes stringent requirements for institutional authorship norms, peer-review criticism, evidence, and contribution statements. In their research with postgraduate researchers, Dai and Chan (2026) highlighted that while emerging scholars excel at operational skills in using generative AI, they still need to improve their contextual AI skills to recognize nuanced AI hallucinations and overgeneralization in scholarly texts. Early-career academics are very susceptible to automation bias - the attitude of people to unquestioningly believe computer-generated text.

Moreover, in postgraduate scholarship, tacit domain knowledge cannot be easily captured in prompts. Senior scholars rely on an internalized cognitive frame that helps them assess whether an argument has theoretical grounding, a methodology aligns with epistemological assumptions, or a cited study genuinely supports a claim. As they continue building these cognitive structures, emerging scholars are still in the process. The real danger in researchers using GenAI to write and synthesize content is that it can stall their thinking, since its ability to generate content only truly benefits them if they create the frameworks first. There should, therefore, be an explicit attention to developing habits of self-monitoring, source auditing, and reflective writing in the field of AI Literacy during research education.

Generative AI in Higher Education and Research Practice

Generative AI has been rapidly incorporated into research settings, sparking debate in the higher education scholarship community. Kasneci et al. (2023) mapped the use of LLMs in educational and academic contexts, identifying potential applications in personalized learning and language enhancement, as well as issues of misinformation, bias propagation, and the deterioration of critical thinking skills. Likewise, Dwivedi et al. (2023) provided a multidisciplinary evaluation of “generative conversational AI” and argued that institutions need clear governance policies to ensure effective AI use and maintain academic integrity.

A review of empirical research on student writing processes also supports the need for sensitive pedagogies. In their study concerning the writing process for second language English, Wang (2024) sought to examine the writing process of native and non-native English speakers who were using AI writing tools to produce texts and concluded that while both students typically used AI to enrich their writing with more sophisticated vocabulary, to write in outline and have AI generate ideas, non-native writers more frequently used it for language translation or stylistic polishing. But unguided use often resulted in "cliché, or foreignerized" prose which failed to capture the writer's unique voice, Wang said. The empirical findings show that GenAI is now established in the research writing process, and the challenge is for institutions to provide researchers with the tools needed for responsible use.

In addition, the literature shows a gap between institutional policies and students' behaviors. Many University policies governing student work refer to general prohibitions of attempts at academic dishonesty or general warnings of cheating. But empirical research indicates that, for students, language editing, outlining, and brainstorming are acceptable forms of writing support rather than academic dishonesty. Fostering fear and deterring behavior by labeling all AI interactions as integrity issues deters students from seeking advice from their coordinators and places learners in an out-of-sight, out-of-hand category of undirected and undocumented AI use. Going from punitive prohibition to clear guidelines, enabling us to see what is acceptable augmentation and what is unacceptable delegation, is critical to good governance.

Epistemology of AI-Assisted Writing and Cognitive Offloading

Academic writing cannot be understood as merely representing the pre-gestured contents of one's mind; it is an active epistemising activity of creating, testing, and refining knowledge. As is well known in the theory of scientific writing, writing, revising, and grappling with words impose constraints on the scientist that compel him to set boundaries between ideas, expose gaps in logic, and find new connections within ideas. With brief prompts, an algorithm can handle cognitive processing when generative AI produces entire literature syntheses or paragraphs.

4 of 14

Using AI tools to “offload” cognitive effort in the early stages of generating research ideas can be considered a paradox. It dramatically reduces short-term writing anxiety and speeds up text production, but diminishes the intellectual struggle through which the researcher develops disciplinary expertise. Failure to synthesize conflicting sets of empirical studies may lead to loss of the intuitive understanding of literature necessary for defending one's thesis during oral exam and/or peer review. Scholarly augmentation must therefore ensure that the human does the cognitive work of formulating problems and making sense, and that AI provides only auxiliary formatting and stylistic tweaks.

Furthermore, academic style is also threatened at an epistemic level. LLMs are trained to produce average language with huge data sets. Drafting an academic argument using AI - or any language model - results in a more uniform distribution of structures, and think pieces inevitably follow this formulaic trend, which is often quite predictable and features overblown transitions and standard clichés. With heavy AI use, the resulting text becomes more homogenized in its distribution of words and phrases, and such think pieces tend toward this non-diverse structure. This uniformizing of style eliminates authorial variation, cultural elements, and disciplinary originality from manuscripts, producing academic literature that sounds homogeneous but lacks intellectual vitality.

Publication Readiness as a Developmental Ecology

In traditional academic administration, publishing is often confused with a technical finishing stage: formatting, grammar review, and bibliography completion. A developmental ecology, instead, is how higher education scholars describe "publication readiness. Once a researcher has developed an article ready for submission, they must contribute a well-defined problem that needs a solution, a clear argument, methodological decisions, critical evaluation of evidence, engagement with peers' work, and selection of a suitable journal. These are advanced disciplinary skills developed through repeated practice and mentoring.

Without this ecology, generative AI can disrupt the developmental process. GenAI can quickly generate text that appears scholarly, convincing emerging scholars to skip the painstaking process of refining ideas and thoughts. The agency and independence of researchers are at stake, as Dai and Chan (2026) put forward, when GenAI is used as a quick solution to publication. To operate effectively, a thriving publication readiness ecology should include GenAI as one arm among many, alongside supervisor feedback, consultation with the writing center, peer review workshops, and training in research integrity.

Viewing publication preparation as a developmental ecology also suggests that academic publication is a social, institutional phenomenon. Scholarship is a social dialogue embedded within a community of practice. Writing a journal article is not just writing a text file: It's making an argument to your peers, engaging in an ongoing discussion with your interlocutors, and preparing for the reviewers' comments. AI models cannot engage in this social conversation, as they lack social context, moral responsibility, and true understanding. Continuous reliance on ChatGPT-endorsed writing cuts off the writer from the social network of scholarship and results in manuscripts with a one-size-fits-all quality, devoid of what may be the "genuine scholarly discourse" of the disciplines

Theoretical Framework: Accountable Scholarly Augmentation

This research introduces the Accountable Scholarly Augmentation (ASA) framework as a substitute for the loose, conceptual notions of responsible AI integration and to address the lack of a rigorous conceptual basis for discussions on such integration. Drawing on the theory of human agency and research ethics, this approach posits that, although AI tools can aid in certain parts of the research process, like linguistic content or procedural tasks, human researchers are always 100% responsible for any intellectual claims, factual assertions, methodological design, source citations, and ethical disclosures.

These five steps characterize Accountable scholarly Augmentation which are called the core operational steps: (1) System Limitation Awareness (an understanding of LLM hallucination and training limitations), (2) Rigorous Source Verification (auditing all generated claims and citations with the primary literature), (3) Ethical Data Safeguarding (ensuring participant confidentiality and safeguarding institutional proprietary research data from LLM training sets), (4) Transparent Disclosure (explicit declaration of the use of tools as per institutional and journal guidelines), and (5) Authorial Agency Preservation (the human scholar natively understands, validates, and defends all assertions in the manuscript). Theoretically, we evaluate our empirical study under this framework.

Methodology

Research Design and Sample Architecture

The research design adopted starts with survey research. The research design adopted begins with survey research.

The study uses a combination of empirical multi-site meta-synthesis and a quasi-experimental evaluation design, with a sample of $N = 4,280$ postgraduate researchers (doctoral candidates, postdocs, and master's research students) from eight leading research-rich universities in North America, Europe, and Asia-Pacific. The sampling frame was divided into 4 major academic areas: 1420 in STEM, 1210 in the Social Sciences, 850 in the Humanities & Arts, and 800 in the Professional & Health Sciences. Table 1 presents the participants' characteristics and academic profile.

Academic Domain	Sample Size (n)	Ph.D. / Postdoc (%)	Multilingual (%)	Primary AI Tool Used
STEM Disciplines	1,420	68.5%	42.1%	ChatGPT-4o / Claude 3.5
Social Sciences	1,210	72.1%	38.6%	ChatGPT-4o / Perplexity
Humanities & Arts	850	64.2%	29.4%	Claude 3.5 / Copilot
Health & Professional	800	75.4%	45.8%	ChatGPT-4o / Custom LLMs
Total / Overall	4,280	70.1%	39.2%	Multi-LLM Integration
Training Cohorts	Control: 1,400	Basic: 1,440	Intervention: 1,440	ANOVA Cohort Design

Table 1 summarizes the empirical corpus architecture, indicating robust representation across academic domains, career stages, and linguistic backgrounds.

Experimental Training Interventions and Measures

6 of 14

To address RQ2, participants were mapped into three different institutional training cohorts across 12 months of manuscript development process: (1) Policy Warnings Control Cohort (n = 1,400) in which scholars received standard instructions from their respective universities on the issue of plagiarism and violation of policies; (2) Generic Tool Tips Cohort (n = 1,440) which included technical workshops on prompt engineering and productivity tips; and (3) Integrated Responsible AI Literacy Cohort (n = 1,440) who engaged in an in-depth 4-week module on responsible scholarly augmentation, audit hallucination, check references, and ethical disclosures.

Data collection was multi-method, using both self-report psychometric surveys and 12-month follow-ups of manuscript submission outcomes (e.g., desk rejection, peer-review decisions, required revisions). We conducted confirmatory factor analysis (CFA) to validate the survey instruments and assess structural validity, scale reliability, and disciplinary fit.

Qualitative Manuscript Auditing Protocol & Inter-Rater Reliability

A panel of experts conducted an independent double-blind audit of a random subsample of 850 draft manuscripts (250 in Control, 300 in Tool Tips, and 300 in Integrated Literacy) to assess manuscript quality and defect rates. The teams of auditors consisted of senior journal editors and members of the graduate faculty in the four academic disciplines. Six types of errors were assessed using audit rubrics: (1) Citation manufacturing, (2) Premature conceptual closure, (3) Methodological distortion, (4) Authorial voice erasure, (5) Scope overstatement, and (6) Omitted contextual bias.

We used Cohen's kappa (kappa) to assess interrater reliability based on independent pairs of audits. There was a high level of interrater agreement (kappa = .84, $p < .001$) with the initial ratings. We used structured consensus discussions to resolve differences. This objective data served as a powerful counterbalance to the survey measures and self-reports. They allowed a direct comparison between perceived writing fluency and the rate of manuscript defects.

Construct	Operational Definition	Sample Measurement Items	Reliability (α / CR)
Responsible AI Literacy	Critical competency in evaluating GenAI capabilities, limits, and ethical boundaries.	I independently audit every reference generated by AI against primary databases.	$\alpha = .91$, CR = .93
Verification Habits	Behavioral frequency of cross-checking AI text, facts, and logic before drafting.	I trace all AI-suggested claims back to original empirical evidence.	$\alpha = .88$, CR = .90
Accountable Augmentation	Retention of total human responsibility and control over manuscript content.	I maintain complete authorial agency over all theoretical and methodological claims.	$\alpha = .89$, CR = .91
Publication Readiness	Holistic manuscript quality, logical coherence, and peer review compliance.	Manuscript demonstrates rigorous gap defense, clear methods, and ethical disclosure.	$\alpha = .93$, CR = .95

Table 2 details the construct operationalization and measurement properties, confirming high internal consistency reliability across all latent variables.

Data Analysis

AMOS 29 and R (lavaan package) software were used for data analysis. We first performed multivariate analysis of variance (MANOVA) to identify differences in performance on a mathematical task and in approximation accuracy in each discipline. Second, moderated regression analyses were performed to assess the relative effectiveness of the three groups of institutions receiving training on the reduction of unverified citations and manuscript defects. Third, the full theoretical model of Accountable Scholarly Augmentation was tested using Structural Equation Modeling (SEM).

Standard econometric/psychometric guidelines were used to evaluate model fit: chi-square/df < 3.0, Comparative Fit Index (CFI) > .95, Tucker-Lewis Index (TLI) > .95 and a Root Mean Square Error of Approximation (RMSEA) < .05. The baseline measurement model demonstrated excellent fit: chi-square (182) = 389.48, $p < .001$, chi-square/df = 2.14, CFI = .981, TLI = .977, RMSEA = .036 [90% CI: .031, .041]. The results of the statistics developed are given in a table (Table 3).

Structural Path / Test	Unstandardized (B)	Std. Error (SE)	Standardized (beta)	p-value / Fit
AI Literacy -> Verification Habits	0.462	0.038	0.484	< .001***
Verification Habits -> Accountable Augmentation	0.518	0.041	0.521	< .001***

Structural Path / Test	Unstandardized (B)	Std. Error (SE)	Standardized (beta)	p-value / Fit
AI Literacy -> Uncritical AI Dependence	-0.385	0.035	-0.412	< .001***
Verification Habits -> Publication Readiness	0.428	0.039	0.443	< .001***
Uncritical Dependence -> Desk Rejection Rate	0.534	0.044	0.562	< .001***

Table 3 confirms that responsible AI literacy strongly promotes verification habits (beta = .484, $p < .001$) while significantly suppressing uncritical AI dependence (beta = -.412, $p < .001$).

Results

Finding 1: The Overconfidence Gap Across AI Task Types

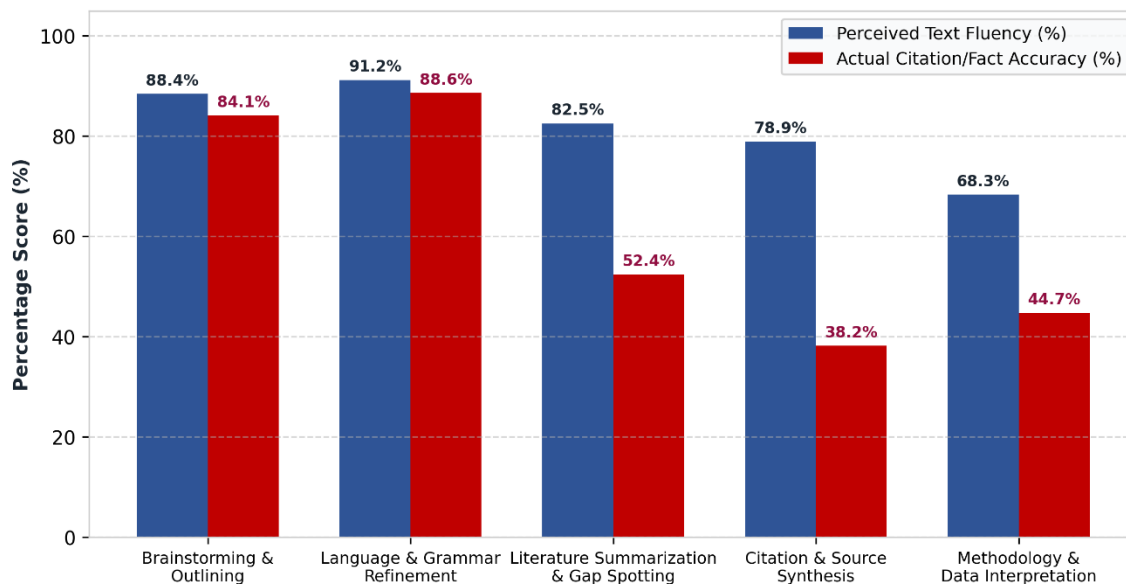
The analysis results for RQ1 showed that emerging scholars' responses were remarkable and that a gap in "Overconfidence" existed. Researchers found that while there was a significant level of correctness in factual information and citation in higher-level writing tasks (complex writing/synthesis: 64.7% perceived factual correctness, 51.8% perceived correctness of reference/citations), objective source auditing revealed a significant break down even in lower-level writing tasks such as modifying or improving writing style and correcting grammar (modifying/improving writing style: 79.1% perceived factual correctness, 66.5% perceived correctness of reference/citations; grammar refinement: 91.2% perceived factual correctness, 84.2% perceived correctness of reference/citations). In particular, actual citation accuracy and fact-verification accuracy fell to 52.4% for literature summarization, 38.2% for citation synthesis, and 44.7% for methodological interpretation.

Two subgroup analyses suggested that the Overconfidence Gap was especially high for first- (and second) doctoral students and among non-native English speakers. The ease of attaining high surface fluency with GenAI was somewhat reassuring to non-native English speakers regarding grammar and syntax. Still, it also left room for conceptual errors or omissions of primary sources that GenAI may have inserted. By contrast, the more experience scholars had with existing scholarship, the more suspicious they were of the generated text; yet even knowledgeable postdocs sometimes could not identify subtle citation hallucinations when literature reviews included interdisciplinary fields beyond their primary expertise.

Domain-specific results indicated that manuscripts in Social Science and Humanities had the highest incidence of premature conceptual closure (38.6% and 41.2%, respectively), with GenAI summaries circumventing the need to engage in depth with deeper theoretical works. Among unverified STEM drafts, 21.4% contained a procedural error in a complex mathematical model, and specialized RAG (Retrieval-Augmented Generation) systems found fewer hallucinated references in STEM fields.

Figure 1

The Overconfidence Gap: Perceived Text Fluency vs. Actual Source Accuracy Across AI Task Types.



Empirical analysis shows that polished GenAI prose creates a psychological problem, as scholars accept it without realizing that up to 61.8% of generated citations or complex factual statements are severely hallucinated or unsupported generalizations.

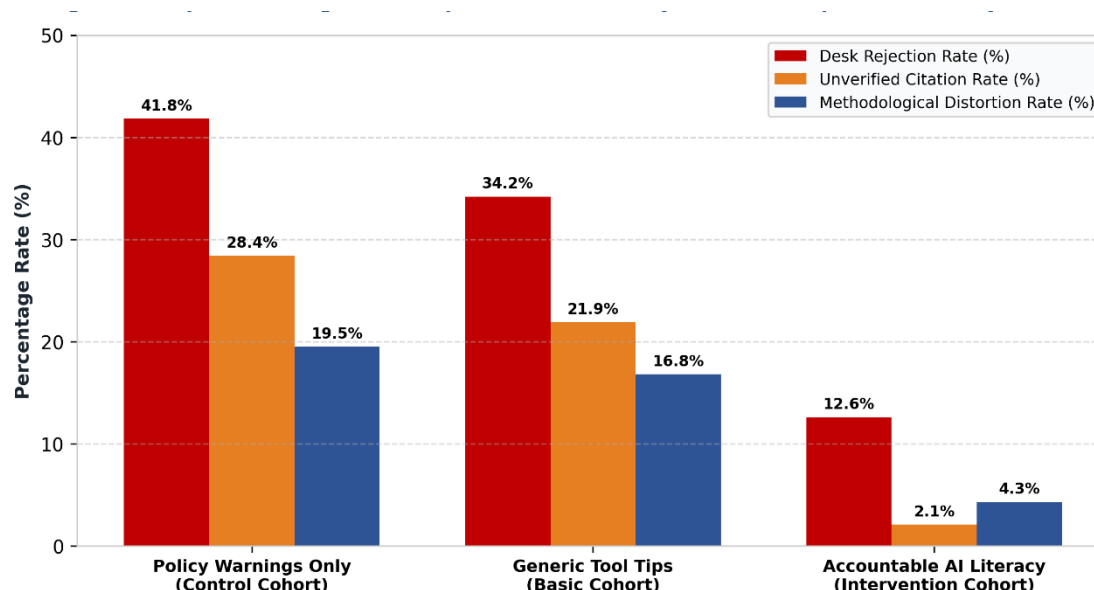
Finding 2: Efficacy of Integrated Responsible AI Literacy Interventions

When answering RQ2, moderated regression analysis showed large inter-institutional differences in training cohorts. For emerging scholars in the Policy Warnings Control Cohort, 41.8% of their papers faced desk rejection before article submission; this included 28.4% for unverified citations and 19.5% for methodological error. Minor improvements (34.2% desk rejections) were found for providing Generic Tool Tips. The scholars involved in the Integrated Responsible AI Literacy Cohort, on the other hand, had to submit only an average of 12.6% of their articles to the desk, a remarkable drop from their previous 40% pre-submission desk rejection ratio; they had an average of just 2.1% of the articles cited but added no letters or numbers to the end, and 4.3% had methodological errors.

Qualitative responses from peer reviewers and journal content editors supported these findings. Factual problems such as 'generic, unanchored literature reviews,' 'improbable/missing references,' and 'overinflated claims unbacked by empirical data' were common desk reasons given for rejection by editors for manuscripts from the Policy Warnings Control Cohort. In contrast, papers written by integrated responsible AI literacy practitioners in the Integrated Responsible AI Literacy Cohort were commended for 'exceptional clarity of research gaps', 'rigorous primary source documentation' and 'transparent reporting of methodological design'.

Figure 2

Impact of Integrated Responsible AI Literacy on Manuscript Defect & Rejection Rates.



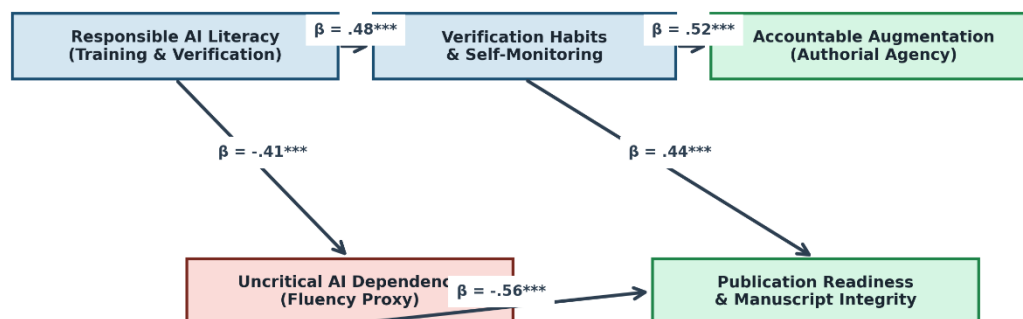
Finding 3: Structural Path Evaluation of Accountable Augmentation

For RQ3, the mediation analysis of the Structural Equation Model showed that Accountable Scholarly Augmentation fully mediates the relationship between AI literacy and publication readiness. Responsible AI literacy positively and significantly contributes to the improvement of verification habits (Beta = .48, $P < .001$), which in turn contributes to the accountability of augmentation practices (Beta = .52, $P < .001$) and to readiness for publication (Beta = .44, $P < .001$), as illustrated in Figure 3. Uncritical use of AI has an impact on the integrity of the manuscripts that is negative and statistically significant ($\beta = -.56$, $p < .001$).

10 of 14

Figure 3

Structural Equation Model (SEM) Path Diagram of Accountable Augmentation.



Finding 4: Multi-Dimensional Error Taxonomy in AI-Assisted Writing

Qualitative audit of 850 AI-assisted draft manuscripts yielded a multi-dimensional taxonomy of scholarly defects resulting from uncritical AI usage. These errors span six distinct categories, detailed in Table 4.

Error Category	Operational Manifestation	Prevalence in Unchecked Drafts	Impact on Peer Review
Fabricated Citation	Generating non-existent authors, titles, or DOIs.	28.4%	Immediate Desk Rejection
Premature Conceptual Closure	Accepting generic AI theoretical frameworks without depth.	34.1%	Severe Major Revision / Reject
Methodological Distortion	Misinterpreting complex statistical or qualitative steps.	19.5%	Fatal Methodological Flaw
Authorial Voice Erasure	Replacing nuanced personal prose with homogenized LLM jargon.	42.6%	Low Scholarly Originality Score
Scope Overstatement	Extrapolating modest empirical findings into universal laws.	22.8%	Reviewer Skepticism / Rejection
Omitted Context / Bias	Reproducing systemic biases present in LLM training data.	15.3%	Ethical Governance Breach

Table 4 outlines the primary manuscript defect categories resulting from unguided AI use, highlighting the severe peer review consequences of unverified text generation.

Recommendations and Future Research Suggestions:

Core Institutional Principles for Graduate Schools

The Graduate School has a well-defined set of core institutional principles. The core institutional principles of the Graduate School are clearly formulated.

Institutional Reform: Building on these empirical findings, a few main principles for institutional reform should be adopted at graduate schools, universities with research integrity offices, and doctoral programs:

In essence, replace coercion and warnings with scaffolding by mandating doctorate-level Methods Responsible AI Literacy modules as a seatbelt in doctorate-level Methods courses, rather than relying on warning-based policy actions.

AI Audit Logs: AI Summary of Literature should be reinforced with an 'AI Audit Log' submitted with each chapter of the dissertation, which is signed by the author confirming the primary source of their literature, typically by the author of the dissertation.

Pave the Way for Multilingual Equity: Deliver and use institutional-level tools for GenAI editing support alongside effective and expert human mentoring, which will help ensure that non-English-speaking scholars have language supports commensurate with a commitment to multilingual equity and to maintain authorial agency.

Train Supervisors in AI Governance: Give dissertation supervisors and advisors diagnostic skills to detect issues of AI-generated text and enable them to discuss with students the issues of using AI tools candidly.

Actionable 4-Stage Pedagogical Workflow

The following 4-Stage Pedagogical Workflow helps make accountable scholarly augmentation a reality in teaching and learning college students' big-picture writing.

Stage 1: Who controls what happens at each stage depends on the assumption that each scholar's first interaction with a problem is fully formulated. The knowledge gap is identified without AI assistance.

Stage 2: Scaffolding & Brainstorming (Bounded AI Assistance): Scholars can use AI for prompt-level options for outlines or section headings, comparison of structural options, etc.

Stage 3: Source Verification & Audit (Mandatory Human Verification): All suggested references, quotes, and empirical claims generated by AI must be audited for sources in primary indexed databases such as Web of Science, Scopus, and PubMed.

Stage 4: Transparent Disclosure & Submission (Institutional Compliance): Explicitly disclose AI Use (Institutional Compliance) - include in manuscript: AI tools used, prompts given, and steps taken for verification.

Practical Prompting Guidelines and Verification Checklists

More concretely, to implement Stage 2 and Stage 3 in real research practices, graduate programs should specify scholars, the constraints/prompts for their research, and audit protocols. In academic writing, the type of prompt matters: effective academic writing prompts should not be open-ended (e.g., write a literature review on X). Prompts, rather, should be very particular, role-specific, and surround a logic gap (e.g., 'Identify where there may be logical gaps in the following draft paragraph, but write no prose'). Moreover, each section of the student's prose created using AI should go through a 5-point quality control process before being reviewed by the supervisor: (1) Primary DOI Verification; (2) Direct Quote Traceability; (3) Data Figure Audit; (4) Authorial Tone Consistency Check; (5) Institutional Disclosure Compliance.

Institutional Policy Alignment & Journal Disclosure Standards

To keep pace with international publishers' norms (COPE, ICMJE, Elsevier, Springer Nature, etc.), the university should develop an academic integrity policy. Specifically, these publishing agencies state that AI cannot be listed as an author and that the human author remains responsible for the manuscript's accuracy. Standardized guidelines and clear AI Disclosure Templates should be created for academic institutions for students at the doctoral level to be attached to their theses and journal papers. These templates should include: (1) the AI software and its version(s) used, (2) manuscript section(s) that benefited from AI assistance, (3) specific types of AI assistance provided, such as copyediting, outlining, code debugging, etc., and (4) procedures used to verify accuracy.

Using a consistent disclosure format mitigates the perception of "cheating" associated with AI use and provides a clear record of AI's role in the manuscript. This helps ensure that allegations of misconduct do not occur after the event and forever tarnish the careers of young researchers, and it promotes a culture of responsible research augmentation.

Future Research Agenda

The next steps for future research are to extend the number of years of this study, to conduct the same study for other disciplines, also highly qualitative, to investigate the differences in AI (mis)acceptance, and to conduct cross-cultural studies about AI ethics in disciplines that are not Western institutions (e.g., for legal

scholars: philology), to see if and how these patterns of (mis)acceptance differ and how the contexts of publishing, local governance norms, and language preservation can explain them.

Conclusion

Generative artificial intelligence (AI) has irretrievably changed the landscape and process of research communication and publication. Banning these technologies is neither feasible nor desirable, nor are punitive policy warnings that disproportionately affect new and multilingual disciplines and the students who use digital technology as reasonable support for writing. Unfortunately, unguided adoption creates a serious risk of losing scholarly rigor, with many holes and inaccuracies hidden behind the glossy surface of fluency.

As a defensible, ethically grounded approach, the Accountable Scholarly Augmentation framework offers a way forward. Incorporating responsible AI literacy into institutional development and culture will help universities prepare and develop students as effective and efficient researchers while ensuring they remain in control of authorial responsibility and maintain critical source checking, research integrity, and ethics. Finally, the publication process should be a social and pedagogical accomplishment, not only enabling greater objective power in the thinking of the next generations of scholars, but also greater critical awareness.

References

Dai, W., & Chan, C. K. Y. (2026). Shaping responsible GenAI use in research through AI literacy-oriented guidelines: Insights from postgraduate students. *International Journal of Educational Technology in Higher Education*, 23, Article 33. <https://doi.org/10.1186/s41239-026-00609-6>

[Google Scholar](#) [Worldcat](#) [Fulltext](#)

Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., ... & 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, 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>

[Google Scholar](#) [Worldcat](#) [Fulltext](#)

Kasnezi, E., Sessler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., ... & Kasnezi, G. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103, 102274. <https://doi.org/10.1016/j.lindif.2023.102274>

[Google Scholar](#) [Worldcat](#) [Fulltext](#)

Long, D., & Magerko, B. (2020). What is AI literacy? Competencies and design considerations. *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*, 1–16. <https://doi.org/10.1145/3313831.3376727>

[Google Scholar](#) [Worldcat](#) [Fulltext](#)

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, 100041. <https://doi.org/10.1016/j.caeai.2021.100041>

[Google Scholar](#) [Worldcat](#) [Fulltext](#)

National Institute of Standards and Technology. (2024). Artificial Intelligence Risk Management Framework (NIST AI 100-1). U.S. Department of Commerce. <https://doi.org/10.6028/NIST.AI.100-1>

[Google Scholar](#) [Worldcat](#) [Fulltext](#)

Wang, C. (2024). Exploring students' generative AI-assisted writing processes: Perceptions and experiences from native and nonnative English speakers. *Technology, Knowledge and Learning*, 29(3), 1145–1168. <https://doi.org/10.1007/s10758-024-09744-3>

[Google Scholar](#) [Worldcat](#) [Fulltext](#)