



Journal of Advanced Research in Computing and Applications

Journal homepage:
<https://karyailham.com.my/index.php/arca/index>
ISSN: 2462-1927



A Framework for Human-AI Synergy: Reinterpretation of the Classical Tripartite Model

Yuslina Yusoff^{1,*}, Wan Nor Hazimah Wan Azib¹, Mimi Zazira Hashim¹, Khadijah Abdul Rahman²,
Fadhilah Mohd Ishak¹, Nur Shaliza Sapiai²

¹ Faculty of Business and Management, Universiti Teknologi MARA Cawangan Kelantan, Machang, Kelantan, Malaysia

² Faculty of Information Science, Universiti Teknologi MARA Cawangan Kelantan, Machang, Kelantan, Malaysia

ARTICLE INFO

Article history:

Received 4 July 2025

Received in revised form 10 September 2025

Accepted 30 September 2025

Available online 6 October 2025

Keywords:

Artificial Intelligence; accountability;
writing assistant; human-AI collaboration

ABSTRACT

This conceptual paper advances a novel framework for understanding the limitations of Artificial Intelligence (AI) writing assistants in academic contexts. Building on prior narrative reviews, we propose a tripartite model categorizing AI limitations into technical, cognitive, and ethical dimensions. Through synthesizing recent literature (2020–2024), we argue that AI's inability to meet academic standards stems from inherent gaps in contextual adaptability, synthetic reasoning, and ethical accountability. Our framework emphasizes the necessity of human-AI collaboration to mitigate these limitations, offering actionable recommendations for scholars, developers, and policymakers. This paper contributes to the discourse on AI in academia by redefining limitations as opportunities for symbiotic innovation rather than mere technological shortcomings.

1. Introduction

The rapid development of AI writing tools like ChatGPT and Grammarly has streamlined academic writing by enhancing drafting, editing, and citation processes. Nevertheless, their conceptual limitations have not yet been thoroughly investigated by researchers. Multiple individuals have examined technical limitations; for instance, [1] identified issues such as inaccurate citations and formatting errors, while [2] emphasized cognitive limitations, specifically AI's inability to process and integrate interdisciplinary concepts or perform critical analyses of written material. Despite the insights from previous studies, a deficiency persists in the continuity that would underpin a conceptual framework for addressing AI's limits as an integral component of a broader systemic issue. Technical limitations, such as static training data, exacerbated cognitive deficiencies, subsequently leading to ethical concerns over plagiarism. This study presents a tripartite model as a foundation for developing a more integrated framework to analyze AI limits concerning technological, cognitive, and ethical dimensions. This approach integrates insights from computer science, cognitive psychology,

* Corresponding author.

E-mail address: yuslinayusoff@uitm.edu.my

<https://doi.org/10.37934/arca.40.1.7685>

and academic moral philosophy, asserting that AI cannot function as a mere tool in academia. Indeed, AI should function as a collaborator, necessitating human brains in the composition and refinement of academic writing.

In order to reduce the possibility of plagiarism and cultural bias, academic ethics should also clarify the moral guidelines[3]. Additionally, this methodology necessitates modifications and extensive methodology. Policymakers must set accountability norms, institutions must teach academics AI literacy, and developers must be transparent about their training data. AI can only go beyond the limits that are now set for it. Instead, diminishing human intelligence and integrity, technology should enhance human cognition and integrity if these endeavors are appropriately integrated into an academic discipline.

2. Problem Statement

The integration of Artificial Intelligence (AI) writing tools into academia has introduced a novel approach to the composing, refining, and formatting of scholarly information. The use of AI in academic research has led to several controversies and difficulties that compromise the integrity of research quality. AI technologies are currently struggling to aid scholarly articles, particularly in conducting reviews, obtaining literature reviews, managing citations, and complying to the unique rules of academic publications [4]. [5] discovered that AI-generated articles can fail to comply with methodological criteria, resulting in outputs unacceptable for peer-reviewed publications. The information offered by AI systems is inadequate for the advancement of interdisciplinary understanding, since AI often relies on produced training data. This constraint leads to superficial assessments, particularly in areas that need significant subject knowledge, such as clinical research or theoretical physics.

Cognitively, AI systems are incapable of maintaining a coherent connection between the diverse fields of knowledge. According to [6], when AI tools integrate concepts from many academic settings, such as sociological theories and economic models, they often result in fragmented narratives. The AI outcomes lack sense, context, and the capacity to comprehend subjective cultural or situational elements. This is especially challenging in fields such as anthropology or philosophy, which rely heavily on context and interpretation for meaningful academic debate. The ethical foundation of AI is the collective human knowledge or pre-existing data patterns [7]. Consequently, it prompts ethical integrity concerns and challenges about uniqueness and inclusion. [8] indicated that writings generated by AI get their wording from the training data.

As a result, there is a risk of employing adjacent phrases, whether deliberately or inadvertently, leading to normative paraphrasing that may result in unintentional plagiarism. AI faces challenges in addressing diverse writing styles, cultural norms, and interdisciplinary frameworks. An example is the AI-generated literature review of existing work on Indigenous practices, which may elevate literature globally while neglecting regional scholars, thereby contributing to epistemic colonialism [9]. The integration of technical, cognitive, and ethical limitations may present challenges to the credibility, inclusivity, and research integrity of AI-assisted outputs. Despite advancements in natural language processing, AI writing tools cannot adequately substitute for human scholars, particularly in tasks requiring critical thought, creativity, or ethical nuance[10]. This study addresses the existing gap by providing a framework for scholars to engage with AI technology, facilitating innovation while upholding scholarly rigor, cultural awareness, and ethical accountability.

3. Research Objectives

1. To examine the impact of technical limitations of AI writing tools on academic writing outcomes
2. To examine the impact of cognitive limitations of AI writing tools on academic writing outcomes
3. To examine the impact of ethical limitations of AI writing tools on academic writing outcomes

4.0 Research Questions

1. What is the impact of the technical limitations of AI writing tools on academic writing outcomes?
2. What is the impact of the cognitive limitations of AI writing tools on academic writing outcomes?
3. What is the impact of the ethical limitations of AI writing tools on academic writing outcomes?

5.0 Hypotheses

1. There is a significant relationship between technical limitations in AI writing tools and the quality of academic writing outputs.
2. There is a significant relationship between cognitive limitations in AI writing tools and the quality of academic writing outputs.
3. There is a significant relationship between ethical limitations in AI writing tools and the quality of academic writing outputs.

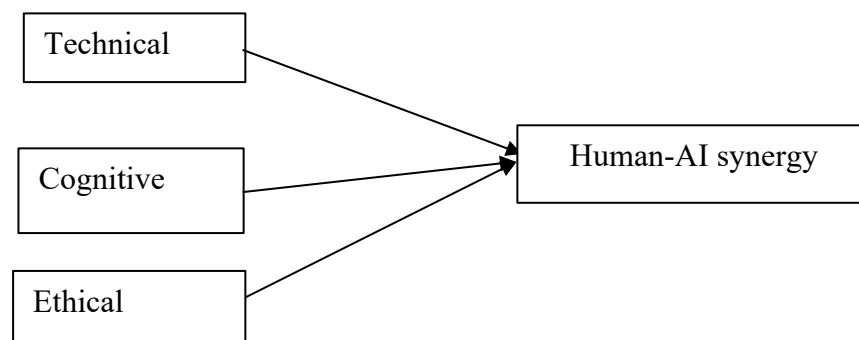


Fig. 1. Conceptual Framework: Tripartite Model of AI Limitations in Academic Writing

5. Conceptual Analysis

The dependence of AI on pre-trained datasets limits its ability to produce context-specific outputs, especially in qualitative research. AI tools often misinterpret theoretical frameworks to qualitative methodologies, including grounded theory and phenomenology [11]. They emphasized that AI systems may combine inductive and deductive approaches, which is problematic for rigorous academic inquiry. This misconception mostly arises because AI models emphasize statistical patterns in language instead of the epistemological precision demanded by qualitative research. This indicates

that they generate outputs resembling academic discourse, but do not align with existing disciplinary paradigms [12].

AI-generated literature reviews in humanities disciplines often give incorrect credit to significant publications or exclude crucial arguments [13][14]. The findings indicate that these errors stem from deficiencies in the datasets utilized for AI training, rather than a deficiency in its analytical capabilities. This restriction highlights a mismatch between the degree of context required for sound scholarly conclusions and the performance of artificial intelligence systems. Due to AI's inability to manage complex discussions, the outputs frequently exhibit linguistic coherence but fail to satisfy the rigorous criteria required for academic work.

[15] further explore these concerns by investigating the degraded performance of pre-trained neural networks when confronted with novel, context-specific datasets. This degradation emphasizes the significance of customized training data in improving the contextual comprehension and usability of AI in a variety of academic fields. The challenges posed by data inadequacy remain prevalent, indicating that it is imperative to continue to improve AI training methodologies to ensure that AI systems can provide pertinent insights without compromising the integrity of scholarly discourse.

Moreover, current research indicates that while AI tools can be beneficial for knowledge management and digital resource management, the efficiency they accomplish does not compensate for the lack of understanding that is essential for superior scholarship. AI interventions in business settings frequently lack insight, since they rely on generalized data sources instead of the contextual backgrounds essential for significant application [15]. It is therefore imperative to provide AI with more sophisticated training data that embodies disciplinary integrity to prevent the propagation of knowledge gaps and misinterpretations that are inherent to its design.

The integration of AI into academic environments necessitates a careful re-evaluation of training methodologies and dataset selection to improve its contextual relevance and accuracy. [16] in their comparative study of AI-generated versus human-edited research drafts, revealing that 79% of AI outputs required substantive revisions to correct factual inaccuracies, contextual misalignments, or citation errors. AI technologies often misinterpret multidisciplinary techniques, such as applying econometric models to ethnographic data without understanding fundamental incompatibility [17].

5.1 Cognitive Limitations: The Human-AI Gap

A distinctive feature of human cognition that is severely absent in AI systems is abductive reasoning, which is the capacity to produce plausible hypotheses from partial or ambiguous facts. GPT-4 and other AI technologies cannot generate novel hypotheses for experimental psychology research. According to an analysis, AI systems is lack of originality in filling in new gaps [18]. AI's reliance on statistical correlations rather than the complex causal or contextual reasoning necessary for hypothesis generation is the cause of this problem.

Additionally, AI systems often lack the methodological proficiency needed to recognize study defects like validity issues or sample biases. In research assessing AI tools in peer-review settings, AI were unable to spot crucial mistakes such insufficient blinding or incorrect power estimates [19]. However, human reviewers were able to detect these methodological flaws with a remarkable 92% accuracy rate [19]. This notable disparity emphasizes how challenging it is for AI to place methodological decisions in the broader framework of ethical or academic norms. One significant issue is that AI often confuses qualitative and quantitative validation techniques, especially in mixed-methods research designs, which may result in inaccurate or deceptive interpretations of important academic work.

In conclusion, there are major obstacles to AI's successful use in academic work because of its limits in producing novel ideas, identifying methodological errors, and comprehending culturally particular circumstances. In order to better match AI systems with the critical thinking abilities that define human cognition and rigorous academic inquiry, recent research emphasizes the urgent need for advancements in AI systems, such as the inclusion of more complex reasoning frameworks and contextually appropriate training data.

5.2 Ethical Limitations: Originality and Cultural Sensitivity

Written academic content generated by AI may include unplanned plagiarism as it is based on pre-existing text patterns in training data. This happens because AI models prefer statistical likelihood over creativity, and they are known to reiterate from high frequency sources in their training data [20]. Such a behavior falsifies academic integrity conventions and devalues the scholarly integrity of the AI-assisted outputs, unsupervised use of AI tools wearing down confidence in academic responsibility and accountability.

AI writing tools are commonly culturally insensitive, mirroring the Western-centric biases that are baked into the data informing their development. This type of bias is an example of academic colonialism as pointed out by [21] where they found that, with respect to the use of AI to create literature reviews on African socio-economic topics, Western scholars were cited, rather than African based authors. These provide highlights and evidence for the ethical necessity to decolonize AI training datasets in order to nurture an inclusive discourse within academia.

If everybody starts using AI tools this may make it harder to differentiate in academia and could lead to a world where the ideas output from academia are too homogeneous. Such homogenization likely stems from AI's incentive to optimize for "safe" language patterns, which deters innovative, and perhaps unusual, scholarly efforts [22]. According to [23] using AI to write in the academia may lead to an intellectual monoculture that suppresses unique perspectives in favor of algorithmic agreement.

6. Methodology

6.1 Research Design

This study will employ a quantitative descriptive cross-sectional survey design to examine the limitations of AI writing tools in academic writing among UiTM Kelantan lecturers. The second-degree cross-lagged design enables systematic data collection at one point in time to examine associations between constructs (AI limitations and academic writing outcomes) and test hypotheses.

6.2 Target Population and Sampling

The population of this study will be all academic lecturers in Universiti Teknologi MARA (UiTM) Kelantan, Malaysia (N = 250) which is located covers few faculties namely. This population comprises of a diverse group of scholarly academics who are engaged in research, publication, and teaching. Consequently, it provides a solid foundation from which to investigate the influence of AI writing tools on scholarly activities. The present research also applies the quantitative approach which seeks to quantify the data and typically applies some form of statistical analysis method in conducting the research. This enables the researcher to achieve precision by focusing on the numerical aspect of the survey data.

The sampling frame will be limited to educators that have used AI writing tools (such as ChatGPT, Grammarly) to support academic tasks. This criterion helps to guarantee that the participants have gained hands-on experience of AI technologies and correspond to processes and places. To gain proportional distribution among faculties, stratified random sampling may be used. The population is stratified into three strata based on faculty membership. Stratification minimizes sampling bias and guarantees that a proportionate number of responses from lecturers of the different disciplines are represented.

The sample size to obtain a representative sample of 152 lecturers is computed with the formula by Krejcie and Morgan (1970) at 95% confidence level and 5% margin of error. This formula has been known throughout the social science literature as a method for choosing sample sizes that are robust while still in reach. The last one applies the following two inclusion criteria to select the sample: Full-time lecturer and (2) experienced with AI writing tools for academic work (e.g., write a research paper, revising a manuscript, or creating a literature review). These criteria guarantee that the data presented represents the problems and opportunities encountered by contemporary researchers when using AI in their work, thus increasing the study's validity and applicability. The self-administered questionnaire is a quantitative method for collecting the needed information and data from the source directly. The questionnaire designed for this study was originally drafted in the English language.

6.3 Data Analysis

The study may use a diverse statistical method to analyze the data and test the hypotheses. First, descriptive statistics which cover frequencies, percentages, means, and standard deviations are used to summarize the demographic data and responses to the survey questions followed by inferential statistics to explore the correlation between variables. Meanwhile, both *t-tests* and *ANOVA* are suitable to be used to compare differences in perceived AI limitations across different faculties and experience levels. *Multiple regression analysis* is to determine how technical, cognitive, and ethical limitations predict academic writing quality. *Chi-square tests* are used to assess the associations between AI tool usage and plagiarism incidents. Lastly, *SmartPLS 4.0*, a Structural Equation Modeling may be used to test hypothesized relationships within the framework of technical, cognitive, and ethical limitations also to quantify how human-AI collaboration moderates the impact of AI limitations on writing quality.

7. Limitations

While the study was properly designed, there are two potential problems to consider. Firstly, for the lecturers who chose not to participate may differ significantly from those who did, potentially affecting the results, this is sampling bias problem. Secondly, is self-reporting bias, this deals with a concern because participants might overstate positive AI experiences or under report plagiarism incidents due to social expectations. Even stratified sampling and anonymity help mitigate these risks, interpretations of the findings should be made cautiously, to enhance validity, future research could combined survey data with AI-generated text analyses.

Table 1
Research Design Plan

Component	Description	Details	Tools/References
Research Design	Quantitative descriptive cross-sectional survey	Examine AI limitations and academic writing outcomes at a single time point	
Population	UiTM Kelantan lecturers using AI tools	Total population of lecturers. Sampling frame: Lecturers with AI usage in past 12 months.	
Sampling Technique	Stratified random sampling	Strata	Krejcie & Morgan (1970)
Data Collection	Self-administered questionnaire (online survey)	Sections: Demographics, Technical/Cognitive/Ethical Limitations, Human-AI Collaboration.	
Data Analysis	Descriptive and inferential statistics; SEM	Descriptive (mean, SD), t-tests/ANOVA, regression, chi-square, SEM (SmartPLS 4.0).	SmartPLS 4.0, Turnitin (plagiarism),
Ethical Considerations	Informed consent, confidentiality, bias mitigation	Digital consent, anonymized data, stratified sampling to reduce faculty bias.	
Limitations	Potential biases	Sampling bias (non-response), self-reporting bias (social desirability).	

8. Implications of study

Undoubtedly, there are several important AI writing tools that are becoming popular amongst higher learning institutions with no difficulties in completing various types of complex assignments. These AI offers voluminous benefits to its users to maximize the quality of traditional work outputs. Hence, the use of it by scholars and academic community can be well embedded to protect institutions' reputable image to have no doubts at all. The proper way to accomplish this is to be able to ensure academic integrity and honesty with all the outputs of academic works can be both appropriate.

The success of using AI writing tools in academia relies on how we can facilitate some cooperation between scholars, developers, and policymakers to find an equilibrium between efficiency and integrity as they relate to academic work. Scholars must use AI as a tool for drafting, not as a stand-in for human work. For example, using AI writing to draft may be helpful as AI can produce drafts rapidly, but AI often lacks the quoted material, rigor, and originality of intellectual scholarship. AI writing may also have incorrect citations, flawed and simplified arguments that lack depth, or even produce some version of unintended plagiarism, so there will always be a need for rigorous human editing for scholarly work to really be a real human edited version not AI works. Scholars should work together towards using AI as a collaborative drafting tool because it will enable us to discover its efficiencies and freedom of creative and critical thinking, both central tenets of academic work.

Developers, in turn, must embrace transparency in the AI systems they develop. Transparency means that developers disclose where their training data comes from, how decision-making algorithms work, and what biases may be built into their models. [24] mentioned that users (students, for example) can audit and adapt algorithms to comply with disciplinary standards. For example, ChatGPT and a lot of the AI content creation tools available often fail to explain how they sourced their information, leading to outputs that appropriate both past research and misrepresent

marginalized points of view. Given delays in adoption, developers can help build trust and mitigate risks by placing ethical recommendations and explainability on AI tools so the practices of AI tools supplement, rather than undermine, scholarly rigor.

Policymakers have an essential role in developing governing frameworks to regulate AI use in academic contexts. [25] argues for guidelines that impose accountability, such as the requirement to disclose whether something was written with the support of AI and create penalties for submitting AI-generated works that were not edited by students. Policies should also focus on equity gaps, such as whether access to AI tools favors institutions with more resources than others. As an example, equitable global standards for ethical AI training data could work to reduce cultural bias and create more equitable ways of practicing academics.

10. Conclusion

In conclusion, this study reconceptualizes the limitations of AI writing tools—technical, cognitive, and ethical—not as challenges hindering academic practice, but as beneficial constraints that foster innovation. By reconceptualizing these limitations through a tripartite framework, we are championing a paradigm shift towards human-AI synergy, where technological capabilities are aligned with human academic expertise to improve academic integrity and creativity. Technical deficiencies, such as whether AI tools produce bibliographic errors or whether the AI has an out-of-date or static knowledge base, require that we want to see adaptive systems that search current information and also refine their performance as more domain-specific clues create iteratively adaptable systems. Cognitive shortcomings or deficiencies, such as synthesizing ideas across disciplines or rigorous examination of methodology, do signal that scholars will remain central to contextualizing emergent AI products and performance generativity across academic discourse complexity. Ethical considerations, from ideas about plagiarism risk to cultural prudence, require systematic judicial means, like the development of clear curated aspects of training data, to reflect firstly accountability and secondly remediation of academic diversity and integrity.

This study, however, emphasizes the potential of AI in higher education by highlighting the roles of human agency to be both augmented. Various benefits from this mixture of collaboration models will result in positive output; where AI takes away repetitive tasks like formatting and drafting while scholars spend their time in areas of critique, creativity, and oversight as immersed partners would shift the limitations of scholarship into opportunities for innovation. These hybrid models with human feedback in the flow of work of AI that could lessen bias, extend adaptability, and produce culturally relevant outputs. The time is now for policymakers to recognize the value human contributors add to the development of AI products and systems and invite them to co-create responsible guidelines and definitions to guide knowledge that is equitable, transparent, and promotes lifelong learning with artificial intelligence whether that is financial support, or design capacity. When we have done well, this cooperative community will make the work of higher education become at its utmost, by democratizing access to new options of writing, while still retaining what matters most in each academic value. Lastly, by treating AI as a partner in collaborations instead of an alternative; higher education will enter the distinct period of the digital age without compromising its foundational principles.

References

- [1] W. H. Walters and E. I. Wilder, "Fabrication and errors in the bibliographic citations generated by ChatGPT," *Sci. Rep.*, vol. 13, no. 1, pp. 1–8, Dec. 2023, doi: 10.1038/S41598-023-41032-5;SUBJMETA=258,639,705,794;KWRD=INFORMATION+TECHNOLOGY,SOFTWARE.
- [2] S. D. Baum, "Artificial Interdisciplinarity: Artificial Intelligence for Research on Complex Societal Problems," *Philos.*

- Technol.*, vol. 34, no. 1, pp. 45–63, Nov. 2021, doi: 10.1007/S13347-020-00416-5/METRICS.
- [3] H. J. Kim and H. Uysal, "Shifting the discourse of plagiarism and ethics: a cultural opportunity in higher education," *Int. J. Ethics Educ.* 2020 61, vol. 6, no. 1, pp. 163–176, Nov. 2020, doi: 10.1007/S40889-020-00113-Z.
- [4] G. Wagner, R. Lukyanenko, and G. Paré, "Artificial intelligence and the conduct of literature reviews," *J. Inf. Technol.*, vol. 37, no. 2, pp. 209–226, Jun. 2022, doi: 10.1177/02683962211048201/ASSET/57A60C1C-CE32-4844-BFEE-16DC60A61F1F/ASSETS/IMAGES/LARGE/10.1177_02683962211048201-FIG1.JPG.
- [5] S. Ebadi, H. Nejadghanbar, A. R. Salman, and H. Khosravi, "Exploring the Impact of Generative AI on Peer Review: Insights from Journal Reviewers," *J. Acad. Ethics*, pp. 1–15, Feb. 2025, doi: 10.1007/S10805-025-09604-4/TABLES/1.
- [6] K. Tammets and T. Ley, "Integrating AI tools in teacher professional learning: a conceptual model and illustrative case," *Front. Artif. Intell.*, vol. 6, p. 1255089, Dec. 2023, doi: 10.3389/FRAI.2023.1255089/BIBTEX.
- [7] V. Terziyan *et al.*, "Towards ethical evolution: responsible autonomy of artificial intelligence across generations," *AI Ethics* 2025, pp. 1–26, Jun. 2025, doi: 10.1007/S43681-025-00759-9.
- [8] I. Dergaa, K. Chamari, P. Zmijewski, and H. Ben Saad, "From human writing to artificial intelligence generated text: examining the prospects and potential threats of ChatGPT in academic writing," *Biol. Sport*, vol. 40, no. 2, pp. 615–622, 2023, doi: 10.5114/BIOLSPORT.2023.125623.
- [9] Y. Ofosu-Asare, "Cognitive imperialism in artificial intelligence: counteracting bias with indigenous epistemologies," *AI Soc.*, vol. 40, no. 4, pp. 3045–3061, Apr. 2025, doi: 10.1007/S00146-024-02065-0/FIGURES/4.
- [10] K. Harati, "ChatGPT and AI-Powered Writing Tools: Unveiling Risks and Ethical Challenges in Scientific Writing," *J. Rev. Med. Sci.*, vol. 4, no. 1, pp. 1–6, Jan. 2024, doi: 10.22034/JRMS.2024.493520.1031.
- [11] P. Tschisgale, P. Wulff, and M. Kubsch, "Integrating artificial intelligence-based methods into qualitative research in physics education research: A case for computational grounded theory," *Phys. Rev. Phys. Educ. Res.*, vol. 19, no. 2, p. 020123, Jul. 2023, doi: 10.1103/PHYSREVPHYSEDUCRES.19.020123/SUPPLEMENTAL_MATERIAL.PDF.
- [12] P. Radanliev, "Artificial intelligence: reflecting on the past and looking towards the next paradigm shift," *J. Exp. Theor. Artif. Intell.*, vol. 00, no. 00, pp. 1–18, 2024, doi: 10.1080/0952813X.2024.2323042.
- [13] M. Mostafapour, J. H. Fortier, K. Pacheco, H. Murray, and G. Garber, "ChatGPT vs. Scholars: A comparative examination of literature reviews conducting by humans and AI," *JMIR Artif. Intell.*, vol. 3, 2024, doi: 10.2196/preprints.56537.
- [14] J. Zybaczynska *et al.*, "Artificial Intelligence–Generated Scientific Literature: A Critical Appraisal," *J. Allergy Clin. Immunol. Pract.*, vol. 12, no. 1, pp. 106–110, Jan. 2024, doi: 10.1016/J.JAIP.2023.10.010.
- [15] Z. Yang, W. Xu, L. Liang, Y. Cui, Z. Qin, and M. Debbah, "On privacy, security, and trustworthiness in distributed wireless large AI models," *Sci. China Inf. Sci.*, vol. 68, no. 7, pp. 1–15, 2025, doi: 10.1007/s11432-024-4465-3.
- [16] A. Patel and J. Cheung, "Artificial intelligence in sleep medicine: assessing the diagnostic precision of ChatGPT-4," *J. Clin. Sleep Med.*, Apr. 2025, doi: 10.5664/JCSM.11732.
- [17] M. Celestin, P. Johnson, A. A. Dinesh, K. M. Vasuki, and P. J. Asamoah, "RESEARCH PRACTICES", doi: 10.5281/zenodo.14559516.
- [18] F. Mazzi, B. L. School, and E. Jaques Bldg, "Authorship in artificial intelligence-generated works: Exploring originality in text prompts and artificial intelligence outputs through philosophical foundations of copyright and collage protection," *J. World Intellect. Prop.*, vol. 27, no. 3, pp. 410–427, Nov. 2024, doi: 10.1111/JWIP.12310.
- [19] Á. A. Cabrera, A. J. Druck, J. I. Hong, and A. Perer, "Discovering and Validating AI Errors with Crowdsourced Failure Reports," *Proc. ACM Human-Computer Interact.*, vol. 5, no. CSCW2, p. 22, Oct. 2021, doi: 10.1145/3479569/SUPPL_FILE/V5CSCW425VF.MP4.
- [20] W. H. Clark, S. Hauser, W. C. Headley, and A. J. Michaels, "Training data augmentation for deep learning radio frequency systems," *J. Def. Model. Simul.*, vol. 18, no. 3, pp. 217–237, Jul. 2021, doi: 10.1177/1548512921991245;PAGE:STRING:ARTICLE/CHAPTER.
- [21] L. J. Janse Van Rensburg, "Harnessing Artificial Intelligence for Economic Growth in Third-World Countries: A Systematic Review and Bibliometric Analysis".
- [22] R. Watermeyer, L. Phipps, D. Lanclos, and C. Knight, "Generative AI and the Automating of Academia," *Postdigital Sci. Educ.*, vol. 6, no. 2, pp. 446–466, Jun. 2024, doi: 10.1007/S42438-023-00440-6/TABLES/2.
- [23] S. Fazelpour and W. Fleisher, "The Value of Disagreement in AI Design, Evaluation, and Alignment," pp. 2138–2150, Jun. 2025, doi: 10.1145/3715275.3732146.
- [24] Y. Cao, Y. Liu, J. Lai, Y. Cao, Y. Liu, and J. Lai, "Leveraging Artificial Intelligence in Outcome-Based Education: A Case Study of Undergraduate Auditing Curriculum," *Adv. Appl. Sociol.*, vol. 15, no. 2, pp. 60–74, Feb. 2025, doi: 10.4236/AASOCI.2025.152004.
- [25] P. G. R. de Almeida, C. D. dos Santos, and J. S. Farias, "Artificial Intelligence Regulation: a framework for governance," *Ethics Inf. Technol.*, vol. 23, no. 3, pp. 505–525, Sep. 2021, doi: 10.1007/S10676-021-09593-Z/METRICS.

