



JOURNAL OF DIGITAL LEARNING AND DISTANCE EDUCATION (JDLDE)

ISSN: 2964-6685

<https://www.rjupublisher.com/ojs/index.php/JDLDE>



Demystifying AI Hallucinations in Philippine Education: A Literature Review

Lexus Dominic Espiritu TAMANO*, Eliz Gueneverre Baltazar GUINTU, Anne Jamille Flores CARLOS

Gordon College, College of Education, Arts and Sciences, English Department, Olongapo City, Philippines.

*Corresponding author: lexxusespiritu01@gmail.com

Abstract

This literature review examines AI hallucinations in the context of Philippine education. Drawing from studies published between 2023 and 2026, it presents research-based insights into how AI-generated inaccuracies influence teaching, learning, and academic practices. The review explores three major themes: AI hallucinations and Philippine education, academic impacts and risks, and pedagogical implications. It highlights concerns regarding misinformation, fabricated references, inaccurate content generation, and the challenges these pose to educational quality and academic integrity. Overall, this paper aims to provide educators, students, and researchers with a clearer understanding of the opportunities and challenges associated with the growing use of artificial intelligence in educational settings. The findings suggest that while AI tools can enhance accessibility, efficiency, and learning support, their effectiveness depends on users' ability to critically evaluate AI-generated outputs and use them responsibly. Ultimately, this review advocates for a balanced and informed integration of artificial intelligence in Philippine education through the promotion of AI literacy, critical thinking, and ethical technology use.

Keywords: AI Hallucinations; Philippine Education; Critical Literature Review

Acknowledgments: The authors would like to express their gratitude to the educators, researchers, and academic institutions in the Philippines whose pioneering studies and insights on artificial intelligence made this literature review possible. Special thanks are also extended to the colleagues and peer reviewers who provided valuable feedback, suggestions, and technical guidance to improve the structure and clarity of this manuscript.

For citation:

Tamaño, L. D. E., Guintu, E. G. B., & Carlos, A. J. F. (2026). Demystifying AI Hallucinations in Philippine Education: A Literature Review. *Journal of Digital Learning and Distance Education*, 5(1), 2118-2125. <https://doi.org/10.56778/jdlde.v5i1.726>

Introduction

In the global view, artificial intelligence (AI) has become more than just a tool for making one's life easier. It is celebrated as a hyper-competent assistant, yet its greatest flaw is that it is a *brilliant liar*. Evidently, AI is utilized in academic, professional, and even for personal use. This in mind, discrepancies became a trend in the scholar community, where misinformation and disinformation rise despite the vast amount of information available online (Broda & Strömbäck, 2024). Citations and AI generated information are written as-to-like it's made by a human being (Ali Adeeb & Mirhoseini, 2023). This is what we call, AI hallucinations (Del Vicario et al., 2016). It is the phenomenon in which AI systems produce plausible and confident output but is factually incorrect, nonsensical, or ungrounded in the real-world (Ge, 2025). It was also defined by (Santos et al., 2025) as something that generates responses based on patterns in data rather than through genuine understanding or verification of information. Likes wise, (Pantao et al., 2026) noted that “when a model encounters gaps in its training data or misinterprets a complex prompt, it may confidently invent plausible-sounding details, citations, or events that do not exist (Bashar & Nayak, 2025).” This part is an example of how AI works in the scholar community, where bots or chatbots respond to queries without being fact-(Kamel, 2024), which then challenges us to think more deeply about the information that we will be using (Shahrzadi et al., 2024).

Therefore, AI hallucinations, in the context of education can be a negative input or a wonderful creation (Dinç et al., 2024). Generative AI tools might be helpful but it is our responsibility as human beings to think critically and not rely too much or be too dependent on the power of AI. In the Philippines, it was highlighted by (Kuznetsova et al., 2025) that AI can become a powerful ally in shaping a better future for Filipino students (Singh & Namin, 2025). However, realizing this potential requires navigating the severe structural disparities between well-resourced urban universities and marginalized rural or K–12 schools (Nguyen & Welch, 2026), where limited institutional policies and intermediate teacher ICT competence leave classrooms highly vulnerable to unverified (Hadianti et al., 2023), hallucinated outputs (Zhai et al., 2024). This vulnerability is significantly deepened by an existing culture of high technology reliance and weak source-checking (Zhang et al., 2025), which creates conditions that allow fabricated data, misleading academic arguments, and flawed references to quietly seep into everyday student coursework (Chen et al., 2023), learning artifacts, and institutional feedback loops (Ali et al., 2024). Because generative models inherently prioritize statistical probability over guaranteed accuracy (Xu et al., 2025), their deceptive but highly convincing outputs alter how students view these tools, demanding a shift from viewing AI as a flawless, benign assistant to recognizing it as an unpredictable, complex entity (Spiegler et al., 2026).

When local educators and academic leaders proactively restructure their assessment designs to treat these flawed, probabilistic responses not as definitive answers but as pedagogical catalysts for active inquiry, students are forced to shift from passive information consumers to highly critical, skeptical evaluators (Aïmeur et al., 2023). By explicitly embedding active verification protocols, independent textual cross-referencing, and structural skepticism into everyday classroom tasks, the systemic threat of automated misinformation can be successfully transformed into a powerful, live exercise that strengthens independent thought, critical thinking capacity, and ethical academic integrity across local institutions. Ultimately, in a local educational landscape marked by uneven verification behaviors, the presence of AI hallucinations implies that the true value of generative technologies depends entirely on a context-sensitive framework that embeds analytical skepticism and rigorous source-checking into the very fabric of the Filipino classroom.

Methodology

This study employed a qualitative descriptive approach through a Critical Literature Review to examine existing studies, scholarly articles, reports, and credible online publications related to AI hallucinations in education, particularly within the Philippine context. A Critical Literature Review

critically evaluates existing research to identify strengths, weaknesses, assumptions, and gaps. Related literature was gathered from academic databases and online sources such as Google Scholar, ResearchGate, and peer-reviewed journals. The researchers selected studies published between 2023 and 2026 to ensure that the discussions and findings reflected recent developments in generative artificial intelligence and its growing influence in education. Priority was given to studies that discussed AI-generated misinformation, reliability issues, ethical concerns, academic risks, and the pedagogical implications of AI tools in teaching and learning. The collected literature was analyzed using thematic analysis. Through careful comparison of related studies, common patterns and significant ideas were identified and organized into three major themes: (1) AI Hallucinations and Philippine Education, (2) Academic Impacts and Risks, and (3) Pedagogical Implications. These themes structured the discussion and interpretation of findings on the broader implications of AI hallucinations in education.

Results and Discussion

AI Hallucinations and Philippine Education

AI hallucinations are best understood as systemic phenomena that emerge from model design, data limitations, and prompting practices, and in the Philippine educational context they have already begun to reshape everyday academic practices and institutional concerns. Studies report that students and educators frequently use generative AI for drafting, research, and assessment tasks, yet verification behaviors remain uneven and often weak, producing conditions in which confidently stated but incorrect outputs can be absorbed into coursework and institutional records (Lindo & Panes, 2024). This pattern of high reliance coupled with low source-checking reflects automation bias and suggests that hallucinations are not isolated model failures but vectors through which misinformation can enter learning artifacts and assessments, thereby undermining the epistemic reliability of student work and teacher feedback (Pedroso et al., 2024). The literature further indicates that these dynamics are amplified by structural constraints in the Philippine system: large disparities in AI readiness between urban higher education and K–12 or rural schools, intermediate levels of teacher ICT competence without specialized GenAI training, and limited institutional policies on AI use all reduce the capacity of schools to detect, correct, or prevent hallucinated outputs from propagating. Crucially, reviews and bibliometric analyses highlight conceptual fragmentation and the absence of standardized definitions and validated detection tools, which impede cross-study comparison and leave stakeholders without common metrics for measuring prevalence or harm. Interpreting these findings in relation to our review, the significance is twofold: first, AI hallucinations function as diagnostic indicators of broader socio-technical vulnerabilities in Philippine education rather than merely technical anomalies; second, because existing evidence skews toward higher education and urban settings, current assessments of risk likely underrepresent the differential impacts on basic education and marginalized communities, creating critical knowledge gaps that must be addressed to form proportionate, context-sensitive policy responses (Ncube et al., 2026).

Academic Impacts and Risks

Studies consistently identify AI hallucinations as a major concern in educational settings because they can affect learning outcomes, academic integrity, and students' ability to evaluate information critically. Hallucinated information may become embedded in students' knowledge structures, resulting in misconceptions that are difficult to correct. Similarly, AI tools like ChatGPT can provide inconsistent and incorrect responses, making it an unreliable source of information when students fail to verify generated content. Later agreed that hallucinations can mislead learners and negatively influence academic outcomes when inaccurate information is accepted as factual knowledge. These findings suggest that although generative AI can support learning, hallucinated

outputs may compromise knowledge acquisition and understanding. The literature also highlights concerns regarding the reliability and trustworthiness of AI-generated content. (Adejumo et al., 2026) argued that hallucinations represent a significant challenge to accuracy and trust in digital learning environments because AI-generated responses often appear convincing despite containing false information. In a similar vein, (Sun et al., 2024) noted that hallucinations create uncertainty among students, educators, and researchers who rely on AI tools for academic tasks. Because generative models prioritize statistical probability over guaranteed accuracy, these deceptive outputs cause profound user frustration and generate pervasive uncertainty, necessitating explicit institutional AI literacy programs to train users in critical evaluation and verification protocols. This systemic unpredictability alters how students view artificial intelligence, shifting their perception of these systems from benign, helpful assistants to unfamiliar and potentially deceptive entities. This shows that while AI technologies offer convenience and accessibility, their educational value depends heavily on the accuracy and verifiability of the information they provide. Several studies emphasize the risks that hallucinations pose to academic integrity. (Wirtz et al., 2026) established a technical taxonomy of AI-generated distorted information, classifying "unfounded fabrication" and "factual errors" as critical issues that directly compromise the objective reliability of text outputs. Likewise, (Li & Yang, 2024) reported that AI systems may generate fake references and misleading academic content that can be unknowingly incorporated into assignments and research papers. (Bittle & El-Gayar, 2025) further demonstrated that AI hallucinations function as critical ethical vulnerabilities that directly influence academic honesty intentions and educational decision-making among undergraduate students. These findings indicate that hallucinations are not merely technical bugs, but ethical concerns that can affect the credibility of academic work and scholarly practices. The reviewed studies also reveal the influence of AI hallucinations on students' critical thinking and dependence on technology. (Yousaf, 2025) warned that repeated exposure to convincing but inaccurate information may reduce critical engagement with learning materials. (Noorbehhahani & Oyibo, 2026) similarly argued that hallucinations can contribute to a decline in critical thinking when students rely excessively on AI-generated responses without verification.

However, exploring general classroom integration, (Deep & Chen, 2025) found that students maintain highly positive attitudes toward generative AI and that its use can actually enhance critical thinking capacity when paired with analytical evaluation. Even when used to assist in technical domains like mathematics or logic, the risk of unverified over-reliance looms large, meaning its true value requires balanced usage alongside traditional instructional resources (Mácha, 2026). Bridging these perspectives, it becomes evident that while AI hallucinations create clear risks for passive, independent thinking, they can also paradoxically serve as a pedagogical catalyst that encourages students to become more careful, skeptical, and rigorous in evaluating digital information. Hence, understanding the academic impacts and risks of AI hallucinations may help educators and students develop more responsible and informed approaches to the use of generative AI in education.

Pedagogical Implications

The systemic integration of generative AI into Philippine education, paired with the persistent risk of model hallucinations, necessitates a fundamental shift from passive content consumption to a process-oriented, "critique-first" pedagogy. Because hallucinations emerge from model design rather than isolated glitches (Osler, 2026) and because these tools prioritize statistical probability over guaranteed factual accuracy, traditional product-oriented assignments face severe ethical and structural vulnerabilities (Desitha Cahya et al., 2026). To prevent fabricated references and misconceptions from embedding themselves into students' knowledge structures, educators must redesign assessments to treat AI outputs not as final answers, but as pedagogical catalysts for critical evaluation (Villarama, 2025). This requires institutionalizing structured AI literacy programs that shift a student's perception of AI from a flawless assistant to an unpredictable entity requiring rigorous verification (Rapp et al., 2025) while ensuring that foundational domains are

balanced alongside traditional instructional resources to mitigate automation bias (Sosas et al., 2026). Crucially, because teacher ICT training is uneven and major structural disparities exist between resourced urban higher education and marginalized, rural K–12 schools (Park & Nan, 2026), these pedagogical adaptations cannot be uniform. Instead, in the absence of standardized tools and metrics, local educational institutions must formulate context-sensitive policies that actively scaffold AI usage (Maleki et al., 2024), transforming the classroom into an environment where rigorous source-checking and analytical skepticism are structurally required to protect the epistemic reliability of learning artifacts and teacher feedback.

Conclusion

The findings of this study demonstrate that artificial intelligence hallucinations are not merely isolated technical anomalies, but are systemic socio-technical vulnerabilities that deeply disrupt the epistemic reliability of Philippine education. The widespread academic reliance on generative AI, coupled with a prevalent automation bias and deficient source-checking behaviors, creates clear pathways for factually incorrect data and fabricated citations to contaminate student coursework, mislead learners, and compromise academic integrity. Furthermore, these risks are heavily magnified by internal structural constraints, specifically the wide disparities in digital readiness between resourced urban universities and marginalized rural or K–12 schools, alongside an intermediate level of teacher ICT training and a lack of standardized diagnostic tools. Ultimately, while hallucinated outputs threaten independent thinking and foster conceptual misconceptions, they can paradoxically serve as a valuable pedagogical catalyst if classrooms actively pivot away from passive information consumption toward rigorous, process-oriented verification protocols. Therefore, to protect academic integrity and educational quality against the systemic threat of automated misinformation, local educational institutions must urgently move away from traditional, product-oriented grading and instead transition toward process-oriented, "critique-first" assessment designs. Academic leaders should develop and institutionalize comprehensive, context-sensitive AI literacy programs that explicitly train both educators and students to view generative models as unpredictable statistical systems requiring constant, manual cross-referencing against verified sources. Finally, because of the severe digital readiness disparities across regions, the Department of Education (DepEd) and the Commission on Higher Education (CHED) must avoid "one-size-fits-all" mandates and instead prioritize specialized generative AI training and defensive, low-tech verification resources for underfunded rural schools and basic education institutions.

Conflicts of Interest

The authors declare no conflict of interest.

Funding

This research received no external funding.

Author contribution

Lexus Dominic Espiritu Tamaño, Eliz Gueneverre Baltazar Guintu, and Anne Jamille Flores Carlos: Conceptualization; Methodology; Formal Analysis; Data Curation; Writing – Original Draft Preparation; Writing – Review & Editing. All authors have read and agreed to the published version of the manuscript.

Declaration of Generative AI in Scientific Writing

The authors utilized generative AI tools during the initial drafting phase for structural outlining and subsequent language refinement. Crucially, all core concepts, data interpretation, and literature analyses were conceived and verified solely by the authors. The authors retain full ownership of and accountability for the manuscript's content and conclusions.

References

- Adejumo, A. A., Oyelere, S. S., Sanusi, I. T., & Suhonen, J. (2026). A systematic review of the impact of GenAI on learning performance, AI hallucinations, and problem-solving in computer science education. *Computers and Education: Artificial Intelligence*, 10, 100570. <https://doi.org/10.1016/j.caeai.2026.100570>
- Aïmeur, E., Amri, S., & Brassard, G. (2023). Fake news, disinformation and misinformation in social media: A review. *Social Network Analysis and Mining*, 13(1), 30. <https://doi.org/10.1007/s13278-023-01028-5>
- Ali Adeeb, R., & Mirhoseini, M. (2023). The Impact of Affect on the Perception of Fake News on Social Media: A Systematic Review. *Social Sciences*, 12(12), 674. <https://doi.org/10.3390/socsci12120674>
- Ali, O., Murray, P. A., Momin, M., Dwivedi, Y. K., & Malik, T. (2024). The effects of artificial intelligence applications in educational settings: Challenges and strategies. *Technological Forecasting and Social Change*, 199, 123076. <https://doi.org/10.1016/j.techfore.2023.123076>
- Bashar, M. A., & Nayak, R. (2025). Chatting with organisational data: A generative AI approach applied to scientific reports for information seeking. *Knowledge and Information Systems*, 67(11), 10657–10690. <https://doi.org/10.1007/s10115-025-02551-x>
- Bittle, K., & El-Gayar, O. (2025). Generative AI and Academic Integrity in Higher Education: A Systematic Review and Research Agenda. *Information*, 16(4), 296. <https://doi.org/10.3390/info16040296>
- Broda, E., & Strömbäck, J. (2024). Misinformation, disinformation, and fake news: Lessons from an interdisciplinary, systematic literature review. *Annals of the International Communication Association*, 48(2), 139–166. <https://doi.org/10.1080/23808985.2024.2323736>
- Chen, S., Xiao, L., & Kumar, A. (2023). Spread of misinformation on social media: What contributes to it and how to combat it. *Computers in Human Behavior*, 141, 107643. <https://doi.org/10.1016/j.chb.2022.107643>
- Deep, P. D., & Chen, Y. (2025). The Role of AI in Academic Writing: Impacts on Writing Skills, Critical Thinking, and Integrity in Higher Education. *Societies*, 15(9), 247. <https://doi.org/10.3390/soc15090247>
- Del Vicario, M., Bessi, A., Zollo, F., Petroni, F., Scala, A., Caldarelli, G., Stanley, H. E., & Quattrociocchi, W. (2016). The spreading of misinformation online. *Proceedings of the National Academy of Sciences*, 113(3), 554–559. <https://doi.org/10.1073/pnas.1517441113>
- Desitha Cahya, Putri Ramdani, Annajmi Rauf, Andi Baso Kaswar, & M Miftach Fakhri. (2026). AI Hallucinations in AI-Assisted Educational Decision-Making and Academic Honesty Intentions Among Undergraduates. *Journal of Applied Artificial Intelligence in Education*, 48–61. <https://doi.org/10.66053/jaaie.v1i2.9>
- Dinç, E., Wherley, M. S., & Sankey, H. (2024). Student Perception of Journaling as an Assessment for an Engagement Experience. *Journal of Experiential Education*, 47(3), 484–503. <https://doi.org/10.1177/10538259231203671>
- Ge, L. (2025). Spectral imaginings and sympoietic creativity: AI hallucinations and the ethics of posthuman creativity. *Big Data & Society*, 12(4), 20539517251400696. <https://doi.org/10.1177/20539517251400696>
- Hadianti, S., Prakoso, T., Brillianing Pratiwi, Memet Casmal, Ucu Rahayu, Siti Aisyah, & Fitra Jaya. (2023). A Compatible Marketing Strategy for the Faculty of Distance Education Institution.

- Journal Of Digital Learning And Distance Education*, 2(5), 589–599. <https://doi.org/10.56778/jdlde.v2i5.161>
- Kamel, H. (2024). Understanding the impact of AI Hallucinations on the university community. *Cybrarians Journal*, (73), 111–134. <https://doi.org/10.70000/cj.2024.73.622>
- Kuznetsova, E., Makhortykh, M., Vziatysheva, V., Stolze, M., Baghumyan, A., & Urman, A. (2025). In generative AI we trust: Can chatbots effectively verify political information? *Journal of Computational Social Science*, 8(1), 15. <https://doi.org/10.1007/s42001-024-00338-8>
- Li, F., & Yang, Y. (2024). Impact of Artificial Intelligence–Generated Content Labels On Perceived Accuracy, Message Credibility, and Sharing Intentions for Misinformation: Web-Based, Randomized, Controlled Experiment. *JMIR Formative Research*, 8, e60024. <https://doi.org/10.2196/60024>
- Lindo, M. R., & Panes, R. (2024). The Hallucination Effect: Correlating Generative AI Usage Frequency with Source Verification Habits among Grade 12 STEM Researchers. *IMCC Journal of Science*, 4(2), 9–11. <https://doi.org/10.65931/d4n7x2h9>
- Mácha, J. (2026). Language Without Propositions: Why Large Language Models Hallucinate. *Philosophies*, 11(2), 42. <https://doi.org/10.3390/philosophies11020042>
- Maleki, N., Padmanabhan, B., & Dutta, K. (2024). AI Hallucinations: A Misnomer Worth Clarifying. *2024 IEEE Conference on Artificial Intelligence (CAI)*, 133–138. <https://doi.org/10.1109/CAI59869.2024.00033>
- Ncube, P. D. N., Dzvapatsva, G. P., Matobobo, C., & Ranga, M. M. (2026). Redefining student assessment in AI-infused learning environments: A systematic review of challenges and strategies for academic integrity. *AI and Ethics*, 6(1), 68. <https://doi.org/10.1007/s43681-025-00871-w>
- Nguyen, D. C., & Welch, C. (2026). Generative Artificial Intelligence in Qualitative Data Analysis: Analyzing—Or Just Chatting? *Organizational Research Methods*, 29(1), 3–39. <https://doi.org/10.1177/10944281251377154>
- Noorbehbahani, F., & Oyibo, K. (2026). AI-overdependence and human cognitive decline: Hazards, evidence, and mitigation strategies. *Computers in Human Behavior Reports*, 22, 101102. <https://doi.org/10.1016/j.chbr.2026.101102>
- Osler, L. (2026). Hallucinating with AI: Distributed Delusions and “AI Psychosis.” *Philosophy & Technology*, 39(1), 30. <https://doi.org/10.1007/s13347-026-01034-3>
- Pantao, J. G., Cahapay, M. B., Ulanday-Lozano, D. M. P., Pelones, M. T. P., Andang-Alaiden, E., Lañojan, R. S., Sales, A. M., & Bagayas, J. C. S. (2026). Examining Artificial Intelligence Literacy of Filipino College Students Using the Theory of Planned Behavior. *European Journal of Educational Research*, 15(3), 699–711. <https://doi.org/10.12973/eu-jer.15.3.699>
- Park, S., & Nan, X. (2026). Generative AI and misinformation: A scoping review of the role of generative AI in the generation, detection, mitigation, and impact of misinformation. *AI & SOCIETY*, 41(2), 1501–1515. <https://doi.org/10.1007/s00146-025-02620-3>
- Pedroso, J. E., Notes, F. A., Ferrer, M. P., Jardeleza, L. C., Estuya, C. B., & Empil, B. D. (2024). Social Media as a Coping Tool Amongst Social Studies Pre-Service Teachers. *JOURNAL OF DIGITAL LEARNING AND DISTANCE EDUCATION*, 3(7), 1154–1167. <https://doi.org/10.56778/jdlde.v3i7.367>
- Rapp, A., Di Lodovico, C., & Di Caro, L. (2025). How do people react to ChatGPT’s unpredictable behavior? Anthropomorphism, uncanniness, and fear of AI: A qualitative study on individuals’ perceptions and understandings of LLMs’ nonsensical hallucinations. *International Journal of Human-Computer Studies*, 198, 103471. <https://doi.org/10.1016/j.ijhcs.2025.103471>
- Santos, A., Cazzamatta, R., & Napolitano, C. J. (2025). Holding platforms accountable in the fight against misinformation: A cross-national analysis of state-established content moderation regulations. *International Communication Gazette*, 87(8), 729–750.

- <https://doi.org/10.1177/17480485251348550>
- Shahrzadi, L., Mansouri, A., Alavi, M., & Shabani, A. (2024). Causes, consequences, and strategies to deal with information overload: A scoping review. *International Journal of Information Management Data Insights*, 4(2), 100261. <https://doi.org/10.1016/j.jjime.2024.100261>
- Singh, S. U., & Namin, A. S. (2025). A survey on chatbots and large language models: Testing and evaluation techniques. *Natural Language Processing Journal*, 10, 100128. <https://doi.org/10.1016/j.nlp.2025.100128>
- Sosas, M. R. D., Fat, A. F. G., Dangculos, J. M. M., Anoba, H. R. P., Remiscal, C. G., & Aleguen, J. M. (2026). Use of ChatGPT as a supplemental learning tool in mathematics. *Contemporary Mathematics and Science Education*, 7(1), ep26004. <https://doi.org/10.30935/conmaths/17989>
- Spiegler, S., Hoda, R., & Pant, A. (2026). Images of AI: How AI practitioners view the impact of Artificial Intelligence on society, now and in the future. *Technology in Society*, 84, 103109. <https://doi.org/10.1016/j.techsoc.2025.103109>
- Sun, Y., Sheng, D., Zhou, Z., & Wu, Y. (2024). AI hallucination: Towards a comprehensive classification of distorted information in artificial intelligence-generated content. *Humanities and Social Sciences Communications*, 11(1), 1278. <https://doi.org/10.1057/s41599-024-03811-x>
- Villarama, J. A. (2025). Demystifying Generative Artificial Intelligence in Academic Classrooms: Students' Attitude and Critical Thinking. *Pakistan Journal of Life and Social Sciences (PJLSS)*, 23(1). <https://doi.org/10.57239/PJLSS-2025-23.1.00314>
- Wirtz, B. W., Weyerer, J. C., & Müller, T. F. (2026). AI-based digital disinformation: A theory-informed and integrated trilateral framework of digital disinformation for information systems research. *International Journal of Information Management*, 89, 103066. <https://doi.org/10.1016/j.ijinfomgt.2026.103066>
- Xu, G., Qian, M., & Meng, L. (2025). Misinformation dissemination on social media: Key research themes and evolutionary paths between 2013 and 2023. *Humanities and Social Sciences Communications*, 12(1), 1775. <https://doi.org/10.1057/s41599-025-06067-1>
- Yousaf, M. N. (2025). Practical Considerations and Ethical Implications of Using Artificial Intelligence in Writing Scientific Manuscripts. *ACG Case Reports Journal*, 12(2), e01629. <https://doi.org/10.14309/crj.0000000000001629>
- Zhai, C., Wibowo, S., & Li, L. D. (2024). The effects of over-reliance on AI dialogue systems on students' cognitive abilities: A systematic review. *Smart Learning Environments*, 11(1), 28. <https://doi.org/10.1186/s40561-024-00316-7>
- Zhang, H., Wu, C., Xie, J., Lyu, Y., Cai, J., & Carroll, J. M. (2025). Harnessing the power of AI in qualitative research: Exploring, using and redesigning ChatGPT. *Computers in Human Behavior: Artificial Humans*, 4, 100144. <https://doi.org/10.1016/j.chbah.2025.100144>