



Publisher homepage: www.universepg.com, eISSN: 2663-7782

British Journal of Arts and Humanities

Journal homepage: www.universepg.com/journal/bjah



OPEN ACCESS | Research Article



Generative AI and Academic Integrity in Higher Education

Sathy Akter* 

Master of Instructional Technology, Touro University of New York, New York, United States of America

*Correspondence: Sornallysathy2017@gmail.co (Sathy Akter, Master of Instructional Technology, Touro University of New York, Manhattan, NY).

Received Date: 18 June 2026

Accepted Date: 19 July 2026

Published Date: 26 July 2026

Abstract

This research report examines the impact of Generative Artificial Intelligence (GenAI) on academic integrity in higher e-ducation on a global scale while specifically analyzing professional education in the USA. In this era marked by the widespread presence of Large Language Models like ChatGPT, DeepSeek and Gemini, the traditional ways of assessment are challenged as never before. This study refers to the literature published in the last five years (2021-2026) to assess the pedagogical, ethical, and technical aspects of the use of AI. The analysis finds that there seems to be an important change in the "honor code" approach of the traditional system for a more complex system that has been recognized as "contract cheating 2.0" and biases of algorithmic detection. Main findings include the opportunities that GenAI brings to personalized learning and student productivity as well as the need to radically reimagine assessment frameworks. The report presents a comparison of the existing different detection methods and draws attention to the potential incompleteness of detection, including the dangers of false positives and bias against non-native speakers. Last, it provides recommendation for both the USA Law Schools and regulators for a future of "AI Literacy" and pedagogical integrity rather than punishment-based surveillance.

Keywords: Generative artificial intelligence, Academic integrity, AI literacy, and Educational technology.

1. Introduction

Research Problem

Generative Artificial Intelligence (GAI) is in the process of revolutionizing education, hence the fidelity of academic integrity. The main problem is distinguishing between AI-generated work and human work, meaning the credibility of both academic institutions and degrees is at serious risk. The current rapid development of LLMs is fundamentally different from the rapid development of technologies in the past, and academic institutions find themselves in a position of having to react to this technology instead of adopting a proactive stance (Mohammadiouno-tikandi and

Babaeitarkami, 2024). The integration of AI technologies in academic work also means that institutions of learning need to reposition their definition of AI-assisted work and shift their focus away from the offensive use of AI. As Gallant *et al.* (2026) observed, students must continue to use AI technologies.

Purpose of the Study

This research intends to analyze the intersection of GenAI with the academic integrity frameworks in higher education. This task aims to overcome the 'ban versus embrace' dichotomy as much as possible and analyze the in-between, how AI, beyond overshadow-

wing the integrity of education, can help in other ways. In the process of constructing a foundation and framework for changes within institutions, the author intends to merge various perspectives and insights, from the technical and practical limitations of proving AI, to the social and ethical concerns of AI bias. The author intends to identify threats that stem from the specific architectural design of AI written texts by analyzing different LLMs. Comparisons will be made between GPT, Gemini, and DeepSeek.

Research Objectives and Hypotheses

- To investigate pedagogical implications of GenAI on student engagement and higher-order thinking skills.
- To assess the effectiveness of existing AI detection systems in terms of fairness.
- To offer a framework for the redesign of assessments, maintaining academic integrity, and recognizing the reality of an AI-enriched professional world.

Research Questions (RQ)

- What is the effect of GenAI tool integration on pedagogical integrity of students in Higher education?

Concepts and Critical Definitions

Table 1: For conceptual clarity, critical definitions used throughout.

Term	Definition	Primary Source
Generative AI (GenAI)	AI systems that can create entities (text, images, code) from patterns that have emerged through training with data.	US Office of Ed Tech (2023)
Pedagogical Integrity	Ensuring alignment between learning objectives, assessment and teaching practices and the creation of effective learning.	Adamakis & Rachiotis (2025)
AI Literacy	Understanding, using, monitoring and critically reflecting on AI applications without compromising on ethics.	Adamakis & Rachiotis (2025)
Contract Cheating 2.0	Using AI to generate technically original academic writing not written by the student.	Gallant <i>et al.</i> (2026)
Student Agency	Students' ability to reflect and act purposefully in the learning process.	Ramezani & Wagner (2025)

Assumptions, Limitations, and Delimitations

Assumptions

The study makes the reasonable assumption that the trend of increasing sophistication in AI reasoning toward more humanlike reasoning will persist. It also views academic honesty as a collective issue, both of institution and student, not necessarily a disciplinary problem.

- How reliable are existing AI detection tools and how do existing AI detection tools bias different student groups?
- What could be the next steps in terms of strategic review and policy adjustment to guarantee academic integrity in the context of AI-enabled work forces?

Significance of Study

The findings of this research are very important for managers and administration of the university, legal education, and legal policy makers. Students are indeed already using GenAI for academic activities; as reported by the Higher Education Policy Institute (HEPI) on a survey of students for 2025 (Freeman, 2025; Banu *et al.*, 2025), more than half of students are using GenAI in their studies, and not always with clear guidance. This study adds value to the field by providing a critical synthesis of the most recent evidence; looking at the requirements of USA Law Schools which focus on precision, attribution and ethical practice. The results are set to help build "AI Literacy" as the core competencies and an essential skill set instead of an optional one.

Limitations and Delimitations

The problem is the field is very changeable, and a 2024 publication might end up being obsolete in 2026. The study is limited to the field of higher education and concentrates on text-based generative models and their effects on academic writing and assessment.

2. Review of Literature

Theoretical Framework: Student Agency and Ethical Responsibility

This study builds on the theoretical underpinning of "Student Agency in the Age of GenAI. In response to this discussion, Ramezani and Wagner, (2025) contend that GenAI does not deprive students of agency, but instead, it redistributes the emphasis of agency from "content production" to "curation and critical oversight. This view is complemented by McCabe's, (2024) decades of study of honor codes that indicate that integrity can be promoted by norms of community instead of surveillance. We close-read these two perspectives. This suggests that the purely technical side of the problem of AI cheating will only be viable when paired with the strong ethical side of the student-teacher relationship. Ramezani and Wagner examine the student's motivation; McCabe thinks of the social contract with the community of scholars.

Opportunities and Challenges of LLMs in Education

Both Zayoud *et al.* (2023) and Kasneci *et al.* (2023) paint a relatively comprehensive picture of the present day context of the 'opportunities and challenges'. Kasneci *et al.* suggest that AI would lead to 'personalised tutoring' and 'improvements in accessibility'. However, there could also be risks of 'cognitive offloading' where students might get more dependent on AI, which could lead to the neglect of developing deep and basic critical thinking skills. This tension has been the focus of Ilgun Dibek *et al.* (2025) in a meta-analysis, which indicates a strong focus on the role of the scaffolding of the thinking skills of AI within a higher order context as offered by the pedagogy of the teachers. The combination of the three presents a continuum of 'pedagogical caution' versus 'techno-optimism'. Whereas Kasneci *et al.* acknowledge that LLMs can have a 'good' impact on the access to education, Zayoud *et al.* are more cautious about the possible 'intellectual atrophy'.

The Impact on Scholarly Communication and Writing

In their 2026 scoping review, Gabay, Funa, and Ricafort highlight that academic writing, particularly among students, is undergoing a significant transition. They describe how Generation AI tools can brains-

form, outline, and refine academic writing. However, the integration of AI in academic writing can lead to a surplus of fabricated citations and the absence of a "human voice" in academic writing (Nassi-Calò, 2025). Nassi-Calò offers a harsh perspective on the state of academic publishing and the resulting cutting of quality among AI-generated "information" and what is published and what can be added to the corpus of knowledge. This stands in contrast to Younas, El-Dakhs, and Noor, (2025) who argue that tools such as DeepSeek and Gemini, when used properly can improve "academic innovation" by enabling researchers to process huge amounts of data in a more efficient manner than manual methods.

Technological Interventions: Detection and Its Discontents

This literature on AI detection shows a highly polarized debate. In their "Heads we win, tails you lose" analysis, Bassett *et al.* (2026) take a more cynical view, suggesting that users of AI detectors may have a false sense of security, while students face pressure for being wrong. Empirically this is backed by Dalalah and Dalalah, (2023) who recorded high number of "false positive" and "false negative". Another major issue that Jiang *et al.* (2024) highlighted is the "bias against non-native English speakers", that, because of their following the standard, less "creative", syntax for learning English language, the patterns in ESL students' written pieces are often mistaken for the characteristics of AI-generated texts. Together this body of work may indicate that the current detection technologies might be doing more harm than good, leading to a 'culture of suspicion', where the issues involve mostly international student populations.

Institutional Leadership and Strategic Planning

According to Khairullah *et al.* (2025), it is responsibility to manage the strategic direction of the use of AI for integration. They compare this with reactive policies, or "emergency-driven policies", during the early days of ChatGPT's release. This is backed by Abedin *et al.* (2026) who argue for the "systematic pedagogical use" of AI, which goes beyond a "culture of suspicion" towards a culture of "collaborative innovation. This needs to be a holistic approach on the institution in the management

structure as well as classroom teaching. The views are synthesized to imply that leadership is about creating a resilient academic culture, one that can deal with rapid technological changes, and one that is "proactive rather than punitive".

Synthesis: From Surveillance to Pedagogical Integrity

The literature review indicates a shift from "detection

based" integrity models to "design based" integrity models. This can be summarized as "the state-of-the-art overview" presented in Adamakis and Rachiotis, (2025) in which AI literacy is a part of the curriculum. According to the literature, cheating prevention is best achieved by rather changing the nature of the "academic task" (Gonsalves, 2025) than by erecting better walls (detectors).

Table 2: Points out the pedagogical shift that was found in the literature.

Aspect	Traditional Model	GenAI-Integrated Model
Focus	Product-based (the essay)	Process-based (the journey)
Strategy	Surveillance and Detection	Literacy and Ethical Agency
Goal	Minimizing Cheating	Maximizing Authentic Learning
View of AI	A threat to be managed	A tool to be mastered

3. Methodology

Research Design

This study uses a systematic literature review (SLR) methodology as outlined by Creswell and Creswell, (2017). The qualitative ‘meta-synthesis’ approach is employed to synthesize the findings presented in multiple studies to gain insight into ‘why’ and ‘how’ AI influences integrity. This approach is selected to allow for sharing of and reflection on divergent viewpoints e.g., technical effectiveness of detector vs pedagogy of AI tools. The study aims to provide a critical lens of the chosen literature to bring topics and themes to the surface that may have been missed in a straight narrative review.

Search Strategy and Sampling

To ensure transparency and replicability, the sampling process was based on Prisma, (2020) guidelines (Page *et al.*, 2021). The databases Scopus, Web of Science and Frontiers in Education were used for the search. Authoritative, peer-reviewed literature from 2021 to 2026 (inclusive) and important policy literature were used in the selection. This timeframe allows the study to consider the immediate pre-ChatGPT and the incredibly rapid development of GenAI.

Inclusion/Exclusion Criteria

Table 3: The thirty references featured at the core of this report are the outcomes of careful deliberation which are summarized.

Criterion	Inclusion	Exclusion
Timeframe	Published between 2021 and 2026.	Published prior to 2021.
Scope	Higher Education context.	K-12 or corporate training only.
Focus	Focus on GenAI and Integrity/Assessment.	General AI or hardware-focused AI.
Peer Review	Peer-reviewed journals or major policy reports.	Blog posts or non-verified social media content.

Data Collection

Data collection was completed using a structured ‘evidence extraction’ method. Each reference was assessed based on their viewpoint on AI detection, the AI detection framework they proposed for the re-assessment of evaluations, and the ethical implications of AI detection. Younas *et al.* (2025) provided a benchmarking comparative review for three tools (DeepSeek, GPT, and Gemini). This systematic

approach enabled the drawing of patterns and themes irrespective of field, geography, or context.

Analysis of Data

The primary method of analysis was thematic analysis, as per Braun and Clarke, (2023). This involved coding and the iterative process of theme development. As AI-generated text is used as a research tool in a particular area, a manual method was employed that is

suited for the interpretation of complex qualitative research data and that retains a critical perspective (Strydom and van der Merwe, 2025). This method was particularly beneficial for identifying nuanced socio-linguistic biases that are likely to remain undetected by automated (socio-linguistic) analyses.

Thematic Findings from the Literature Review

Theme 1: The Dual Impact on Pedagogical Integrity (Addressing RQ1)

There are stark differences between the efficiency improvements promised by AI, versus the educational shortcomings from over-relying on AI. Long *et al.* (2026) find that the influence on student engagement is "mediated by teaching methods. AI works great as a "co-pilot" in brainstorming, but it comes as a problem when it is used as a "proxy" for thinking. That's an indication that it's not technology's fault but rather that the problem may be the "pedagogical vacuum" in which it's being used.

Yermaganbetova *et al.* (2025) present the quasi-experimental findings of using AI-informed learning platforms to enhance "student performance" in terms of quality of output. But Bhatia *et al.* (2026) warn that this does not necessarily result in better "critical thinking. They contend that even a well-written essay on the same topic produced with the help of AI could still indicate that a student does not have deeper conceptual understanding but merely knows how to get the AI to produce that essay. The results indicate that the learning objectives should be equally clear in their alignment with the use of the AI tools, and a

"scaffolding" process for these then should gradually release responsibility for the student's use of the AI tools to the learner, otherwise pedagogical integrity (Adamakis & Rachiotis, 2025) is not maintained.

Theme 2: The Reliability and Bias of Detection (Addressing RQ2)

The most relevant discovery with respect to RQ2 is the "unreliability of algorithmic surveillance". The findings of both (Dalalah and Dalalah, 2023; Bassett *et al.*, 2026) make a strong argument for the transfer of responsibility regarding academic integrity from AI tools to human judgment. It's explained that these instruments are vulnerable to "adversarial attack," which is essentially the process of tweaking AI-generated text to evade detection, but also original human writing.

Moreover, the "two-channel hierarchical attention mechanism" (Chang *et al.*, 2023) and "transformer-stacking frameworks" (Hoque *et al.*, 2025) that have been employed in cutting-edge text classification models reveal that AI can indeed identify AI, but only with a rather large "socio-linguistic bias. The second issue is that, as Jiang *et al.* (2024) stress, the bias against non-native English speakers is an inherent property of the way these detectors work: they "have a prejudice against 'low perplexity' and 'predictable syntax', both of which are characteristics of ESL writing". This puts international students into a "double jeopardy" because they are more likely to be subject to checks for academic integrity anyway.

Table 4: provides an overview of technical hazards to detection found in the literature.

Limitation Type	Description	Key Reference
False Positives	Human work is misidentified as AI-generated.	Dalalah & Dalalah (2023)
Linguistic Bias	Penalization of ESL students' formal syntax.	Jiang <i>et al.</i> (2024)
Evasion	"Humanizing" tools that bypass detection.	Bassett <i>et al.</i> (2026)
Low Reliability	Inconsistent results across different LLMs.	Younas <i>et al.</i> (2025)

Theme 3: Redesigning Assessment for the AI Era (Addressing RQ3)

In response to the integrity gap, the literature overwhelmingly indicates that there should be a shift toward "Contextual Assessment Design" (Gonsalves, 2025). According to Chuang and Yan, (2025) "Language assessment" needs to progress from

focusing on the final product to the "process of writing. This involves the use of oral vivas, in class supervised writing, as well as "AI-reflection logs", in which students record their communication with GenAI tools. Even if such AI is used, this "process-oriented" approach requires the student to show how

and why the content was created, and to show understanding of the process.

However, a global perspective is offered by Cabrera Vélez *et al.* (2026) which states that in South American higher education the emphasis is on "assurance of integrity" through community-based learning, rather than technical issues. This corresponds to the model of "AI-permissibility" rubrics, discipline-specific and university-wide, proposed by Khairullah *et al.* (2025) under the umbrella of "responsible strategic leadership". For example, a creative writing class might require the exclusion of AI, but a computer science class may require that the AI be used for code optimization. Unlike a 'blanket ban' approach which

failed to work, this approach helps to have a more nuanced understanding about what is happening.

4. Results and Discussion

Summary of Thematic Findings

Through thirty thematic analysis sources, it is revealed that the thematic area of the 'AI crisis in education' is a 'crisis of assessment' in education. The research concluded that AI detection tools, despite their technological sophistication (for example, the "transformer-stacking framework" (Hoque *et al.*, 2025), have the following limitations: sociologically faulty and ethically problematic. Switching from 'product' to 'process' is not just a move but an issue of structure the only way to sustain the value of the degree.

Table 5: A brief overview of the main findings based on the research questions.

Research Question	Key Finding	Primary Evidence
RQ1: Pedagogical Impact	AI support through "scaffolding" enhances engagement, while "cognitive offloading" potentially stifles critical thinking.	Long <i>et al.</i> (2026); Bhatia <i>et al.</i> (2026)
RQ2: Detection & Bias	Detectors are not reliable and systematically disadvantage for people whose primary language is not English.	Jiang <i>et al.</i> (2024); Bassett <i>et al.</i> (2026)
RQ3: Assessment Redesign	Change from "output-based," to "process-based" and "contextual" assessment.	Gonsalves (2025); Chuang & Yan (2025)

Research Reflections

The sector is in a transitional period, reflected in the study. To sum up, the scenario for higher education in the future is not "AI-free" but rather "AI-negotiated," as per (Kasneci *et al.*, 2023; Gallant *et al.*, 2026). The most important, or pivotal, observation is that this is not a "technical problem" of AI cheating, but in fact a "pedagogical problem" of assessment design. The institutions were the ones who think of GenAI as a "cheating machine" will not be able to prepare their students for the environment where people think of it as an "efficiency machine." Based on the synthesis of literature, the human element or the unique worldview of students is the starting point to be considered for all pedagogical strategies in the future.

Advice to Regulators and USA Law Schools

The challenge for USA Law Schools is that legal practice is undergoing significant transformation due to GenAI for document review, discovery and drafting. In this context, regulators like the Office for Students (OfS) and the Solicitors Regulation Authority

(SRA) ought to think about a three-part approach which puts ethical skill above technical checking.

Mandating the requirement for "AI Literacy" should be part of legal ethics training. Incorporating LLMs into legal research comes with its share of hazards, known as "hallucinations" and "data privacy". About extremely precise disciplines, Nassi-Calò, (2025) points to the risk of "automated inaccuracy. Second, don't use AI detectors as only evidence of misconduct. However, because of the "bias against non-native speakers" (p. 980), such use could result in unequal outcomes and potential legal challenges under the Equality Act.

Third, promote "Authentic Assessment" reflective of the law profession. One option is for schools to ask students to critically edit a contract produced by an AI or fact-check a case summary produced by an AI. Content production followed by critical evaluation of outputs illustrates an agency evolution of students, as Ramezani and Wagner, (2025) describe. Law Schools can continue preserving professional integrity as they

embrace technology by helping students become “curators” of AI outputs.

5. Conclusion

Generative AI will be a permanent tool in education. However, as this report shows, there are serious consequences that come with the intersection of emerging AI tools and academic integrity, ranging from “contract cheating 2.0” to systemic algorithmic bias. More technology will not solve these challenges, but a change in pedagogy will. “Responsible strategic leadership” and “contextual assessment design” are two important guiding tools that allow institutions to preserve their degrees' worth in a world that is augmented by AI, while also empowering students. More emphasis should be placed on a “human-in-the-loop” approach, focusing on the student's perception of learning, and especially on achieving their educational goals. It is not about AI and technologies, but about how we inspire the community. This commit is to the integrity of the future University.

6. Acknowledgment

The authors wish to thank the members of the Academic Integrity Advisory Board and the Faculty of Business and Law at the University of Greenwich for their valuable input, critical feedback and support in developing the contextual assessment frameworks which are proposed in this study. This work was partly funded by the Educational Development and Innovation Fund.

7. Conflicts of Interest

The authors state that they have no commercial or financial interest that could be interpreted as a potential conflict of interest.

8. References

- Abedin, M. Z., Hayajneh, A., & Raahemi, B. (2026). Pedagogical Use of Responsible Generative AI in Higher Education; Opportunities and Challenges: A Systematic Literature Review. *AI in Education*, 2(2), 11.
- Adamakis, M., & Rachiotis, T. (2025). Artificial intelligence in higher education: A state-of-the-art overview of pedagogical integrity, artificial intelligence literacy, and policy integration. *Encyclopedia*, 5(4), 180.
- Banu AR, Ainy S, and Huda MN. (2025). Internationalization of english in Bangladeshi higher education: evaluating teachers self-efficacy and understanding, *Asian J. Soc. Sci. Leg. Stud.*, 7(4), 358-365.
<https://doi.org/10.34104/ajssls.025.03580365>
- Bassett, M. A., et al. (2026). Heads we win, tails you lose: AI detectors in education. *J. of Higher Education Policy and Management*, 1–16.
- Bhatia, A., Yadav, P., & Sharma, A. (2026). Impact of generative AI tools on students' productivity and critical thinking. *EDPACS*, 1–11.
- Braun, V., & Clarke, V. (2023). Is thematic analysis used well in health psychology? A critical review of published research, with recommendations for quality practice and reporting. *Health Psychology Review*, 17(4), 695–718.
- Cabrera Vélez, J. P., Llanos García, R. V., & Bonilla Fierro, L. F. (2026). Generative AI and the assurance of academic integrity in the context of South American higher education. *Frontiers in Education*, 11, p. 1796556.
- Chang, G., Hu, S., & Huang, H. (2023). Two-channel hierarchical attention mechanism model for short text classification. *The J. of Supercomputing*, 79(6), 6991–7013.
- Chuang, P. L., & Yan, X. (2025). Language assessment in the era of generative artificial intelligence: Opportunities, challenges, and future directions. *System*, p. 103846.
- Creswell, J. W., & Creswell, J. D. (2017). Research design: Qualitative, quantitative, and mixed methods approaches (5th ed.). *London: SAGE Publications*.
- Dalalah, D., & Dalalah, O. M. (2023). The false positives and false negatives of generative AI detection tools in education and academic research: The case of ChatGPT. *The Inter J. of Management Education*, 21(2), p. 100822.
- Freeman, J. (2025). Student generative AI survey 2025. *London: Higher Education Policy Institute*.
- Gabay, R. A. E., Funa, A. A., & Ricafort, J. D. (2026). Generative Artificial Intelligence (GenAI) for Academic Writing in Higher Education: A Scoping Review of Applications, Challenges,

- and Implications. *Inter J. of Edu in Math., Sci. and Tech.*, **14**(1), 200–232.
- Gallant, T. B., Davis, M., & Khan, Z. R. (2026). Academic Integrity in the Age of AI. Elements in Generative AI in Education. *Cambridge: Cambridge University Press*.
- Gonsalves, C. (2025). Contextual Assessment Design in the Age of Generative AI. *J. of Learning Development in Higher Education*.
- Hoque, M. N., *et al.* (2025). Advancing cyberbullying detection in low-resource languages: a transformer-stacking framework for Bengali. *Frontiers in Artificial Intelligence*, **8**, p. 1679962.
- Ilgun Dibek, M., Sahin Kursad, M., & Erdogan, T. (2025). Influence of artificial intelligence tools on higher order thinking skills: A meta-analysis. *Interactive Learning Environments*, **33**(3), 2216–2238.
- Jiang, Y., Hao, J., & Li, C. (2024). Detecting ChatGPT-generated essays in a large-scale writing assessment: Is there a bias against non-native English speakers? *Computers & Education*, **217**, p. 105070.
- Kasneci, E., *et al.* (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, **103**, p. 102274.
- Khairullah, S. A., *et al.* (2025). Implementing artificial intelligence in academic and administrative processes through responsible strategic leadership in the higher education institutions. *Frontiers in Education*, **10**, p. 1548104.
- Long, D. Y., Wang, S., & Lu, X. T. (2026). Artificial intelligence in higher education: a systematic review of its impact on student engagement and the mediating role of teaching methods. *Frontiers in Education*, **10**, p. 1648661.
- McCabe, D. (2024). Cheating and honor: Lessons from a long-term research project. In *Second Handbook of Academic Integrity* (pp. 599–610). *Springer*.
- Mohammadiounotikandi A., and Babaeitarkami S. (2024). The ethics of artificial intelligence: balancing progress with responsibility, *Int. J. Mat. Math. Sci.*, **6**(2), 30-37.
<https://doi.org/10.34104/ijmms.024.030037>
- Nassi-Calò, L. (2025). The use of Generative Artificial Intelligence in Scholarly Communication. *Revista Latino-Americana de Enfermagem*, **33**, p. e4560.
- Page, M. J., *et al.* (2021). The PRISMA 2020 statement: an updated guideline for reporting systematic reviews. *BMJ*, **372**, p. n71.
- Ramezani, D., & Wagner, M. (2025). Redefining student agency in the age of GenAI. In *EDULEARN25 Proceedings* (pp. 7102–7107). IATED.
- Strydom, C., & van der Merwe, S. (2025). Presenting a manual method for complex qualitative data analysis requiring a human perspective. *Acta Commercii*, **25**(1), p. 1461.
- United States Office of Educational Technology. (2023). Artificial intelligence and the future of teaching and learning: Insights and recommendations. *Washington, DC: US Department of Education*.
- Yermaganbetova, M., *et al.* (2025). Evaluating the impact of an AI-integrated learning platform on student performance: a quasi-experimental study. *Frontiers in Education*, **11**, p. 1792353.
- Younas, M., El-Dakhs, D. A. S., & Noor, U. (2025). The impact of artificial intelligence-based learning tools in academic innovation: a review of DeepSeek, GPT, and Gemini (2020–2025). *Frontiers in Education*, **10**, p. 1689205.
- Zayoud, M., Oueida, S., & Ionescu, S. (2023). Impact of ChatGPT on education: Challenges and opportunities. *Inter Conference of Man and Industrial Engineering*, **11**, 75–82.

Citation: Akter S. (2026). Generative AI and academic integrity in higher education, *Br. J. Arts Humanit.*, **8**(4), 782-789. <https://doi.org/10.34104/bjah.02607820789>

Copyright: © The Author(s), 2026. Published by the UniversePG. This is an **Open Access** article, distributed under the terms of the Creative Commons Attribution 4.0 License (<http://creativecommons.org/licenses/by/4.0/>), which permits un-restricted use, distribution and reproduction in any medium, and provided the original work is properly cited. 