

Original Article

ARTIFICIAL INTELLIGENCE AND ACADEMIC INTEGRITY: RETHINKING ASSESSMENT MODELS IN THE AGE OF GENERATIVE AI

Dr. Isha Joshi ^{1*} , Dr. Sumit Maheshwari ²

¹ Independent AI Consultant (Education) and Assistant Professor, Faculty of Management Studies, Medicaps University, Indore, India

² Dean In-Charge, School of Law and Public Policy, Avantika University, Ujjain, India



ABSTRACT

The research paper focuses on how generative artificial intelligence (GenAI) is disruptive to academic integrity and academic assessment practices in higher education. To deal with the critical lag between technological change and institutional control, the study uses a pragmatic mixed-methodology framework that is organized into three phases. Phase 1 will conduct a thematic document analysis of the academic integrity policies of major international universities (Harvard, Oxford, Melbourne, Delhi) in order to build a comparative legal and policy framework on AI-assisted authorship and disclosure. Phase 2 collects primary empirical evidence based on a structured survey (N = 500) of students and teachers in institutions with Management, Law, and Arts programs in Madhya Pradesh, India (Indore, Ujjain, Bhopal). The data is then thoroughly correlated and descriptively analyzed using IBM SPSS through modified versions of the Technology Acceptance Model (TAM) and Academic Dishonesty scales to visualize perceptions of behavioural intentions towards GenAI and perceived efficacy of plagiarism detectors. Phase 3 compares alternative assessment strategies (viva-based, project-based, and timed case analyses) with authenticity and AI misuse criteria. Summing up these stages, the work proposes the new model of academic integrity governance, stating that in order to create technological disruption-resilient assessment, the tripartite alignment of institutional regulation responses and the basic pedagogic reconstruction are needed. The results can allow the university administrators and curriculum developers to obtain empirical evidence and practical models that would assure pedagogical credibility in the AI-ubiquitous age.

Keywords: Generative AI, Academic Integrity Governance, Plagiarism, Assessment Models, Educational Policy, Mixed-Methods

INTRODUCTION

The ubiquitous application of Generative Artificial Intelligence (GenAI), and in particular of Large Language Models (LLMs) like GPT-4 and Claude 3, has led to a paradigm shift in the operation of higher learning [Baidoo-Anu and Owusu Ansah \(2023\)](#), [Dwivedi et al. \(2023\)](#). GenAI systems, unlike the antecedent technologies, which simply served as information retrieval systems, simulate complex cognitive processes: synthesizing unrelated literature, constructing coherent arguments, and writing contextually sensitive academic texts [Cotton et al. \(2024\)](#), [Bearman et al. \(2023\)](#). Although the tools provide deep democratizing prospects of individualized tutoring and administrative effectiveness, they also jeopardize the axiomatic essence of higher education, the real quantification of cognitive achievement of a student based on asynchronous tests [Lodge et al. \(2023\)](#), [Moorhouse et al. \(2023\)](#).

*Corresponding Author:

Email address: Dr. Isha Joshi (ishajoshi.r@gmail.com), Dr. Sumit Maheshwari (sumit.maheshwari@avantika.edu.in)

Received: 15 July 2026; Accepted: 23 August 2026; Published 18 September 2026

DOI: [10.29121/ShodhAI.v3.i2.2026.105](https://doi.org/10.29121/ShodhAI.v3.i2.2026.105)

Page Number: 62-69

Journal Title: ShodhAI: Journal of Artificial Intelligence

Journal Abbreviation: ShodhAI J. Artif. Intell.

Online ISSN: 3048-9245, Print ISSN: 3108-1940

Publisher: Granthaalayah Publications and Printers, India

Conflict of Interests: The authors declare that they have no competing interests.

Funding: This research received no specific grant from any funding agency in the public, commercial, or not-for-profit sectors.

Authors' Contributions: Each author made an equal contribution to the conception and design of the study. All authors have reviewed and approved the final version of the manuscript for publication.

Transparency: The authors affirm that this manuscript presents an honest, accurate, and transparent account of the study. All essential aspects have been included, and any deviations from the original study plan have been clearly explained. The writing process strictly adhered to established ethical standards.

Copyright: © 2026 The Author(s). This work is licensed under a [Creative Commons Attribution 4.0 International License](https://creativecommons.org/licenses/by/4.0/).

With the license CC-BY, authors retain the copyright, allowing anyone to download, reuse, re-print, modify, distribute, and/or copy their contribution. The work must be properly attributed to its author.

A common crisis in academic integrity which is theorized as AI-assisted unauthorized authorship essentially undermines the traditional concept of contract cheating [Bretag et al. \(2018\)](#), [Foltynek et al. \(2023\)](#). Since GenAI probabilistically generates unique text instead of copying indexed databases, it can easily avoid traditional similarity-checking infrastructures [Weber-Wulff et al. \(2023\)](#), [Myers et al. \(2023\)](#). This technical imbalance places institutions of higher learning (HEIs) in an enforcement vacuum of probabilistic AI detection modules that the recent literature has shown to be both inherently biased and have high false positives [Liang et al. \(2023\)](#), [Pande and Roberts \(2023\)](#).

To navigate this disruption, it is necessary to shift to actionable, systematized architecture beyond descriptive concerns.

Image 1

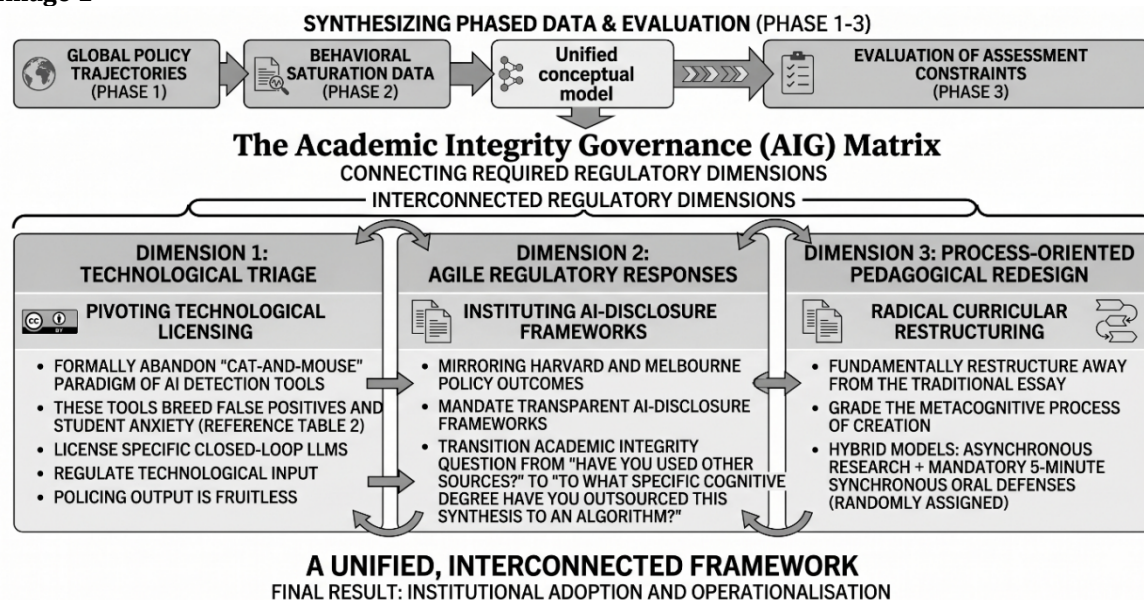


Image 1 Proposed Conceptual Framework for AIG

Source: Author's Compilation

To operationalize this conceptual framework and contribute a scalable model for HEIs, this study is structured across three methodical phases, combining global policy analysis with localized empirical validation.

PHASE 1: DOCUMENT ANALYSIS AND COMPARATIVE POLICY FRAMEWORK

The first layer of the methodology will be a systematic document analysis of academic integrity policies of different, high-ranking academic ecosystems. The goal of this step is to create a comparative background of how elite institutions are institutionally classifying and reacting to the use of AI before attempting to test local empirically crafted perceptions.

SAMPLE AND THEMATIC PROCEDURE:

The publicly available repositories of four benchmark institutions, the University of Oxford, the University of Melbourne, and the University of Harvard University, and the University of Delhi, were searched to extract official academic integrity guidelines, addendums related to the use of AI, and assessment regulatory frameworks as of 2023.

A qualitative thematic analysis was carried out (with NVivo 14 being used in the efficiency of the coding process) to map institutional reactions in four preset regulatory nodes: (a) definitions of AI-generated assignments, (b) the codification of plagiarism versus augmented authorship, (c) compulsory AI disclosure requirements, and (d) institutional requirements of redesigned alternative assessments.

COMPARATIVE POLICY FRAMEWORK

The result of this stage is a clear movement of the absolute prohibition in controlled integration, stratified strongly, with spatial academic cultures.

Table 1

Table 1 Comparative Thematic Analysis of Institutional AI Policies			
Institution	AI-Generated Assignments & Plagiarism	AI Disclosure Requirements	Redesigned Assessment Directives
Harvard University (USA)	Treats unauthorized AI-generated work as academic misconduct; emphasizes student responsibility for originality and intellectual contribution (Harvard University, 2025).	Requires explicit acknowledgment of AI use, including tools, prompts, and integration processes where permitted (Harvard University, 2025).	Encourages course-specific AI policies and controlled assessments (e.g., in-class evaluations) to ensure authenticity (Harvard University, 2025).
University of Oxford (UK)	Treats unauthorized AI use as academic misconduct; emphasizes integrity and independent reasoning (University of Oxford, 2025).	Requires disclosure when AI use is permitted and mandates transparency in academic work (University of Oxford, 2025).	Promotes traditional and controlled assessment formats, including invigilated examinations and oral defenses (University of Oxford, 2025).
University of Melbourne (Australia)	Adopts a regulated integration approach balancing AI use with academic integrity (Jin, 2025).	Emphasizes transparency through formal disclosure mechanisms such as AI usage declarations (Alduais, 2025).	Encourages process-based and reflective assessment aligned with AI-integrated learning (Jin, 2025).
University of Delhi (India)	Regulatory framework is evolving; AI use is broadly governed under existing plagiarism norms (Jin, 2025).	Lacks a standardized institutional disclosure framework; practices vary across departments (Jin, 2025).	Shows gradual shift toward internal, application-based assessments such as presentations and case analyses (Jin, 2025).

Note: GenAI = Generative Artificial Intelligence. The analysis is based on institutional policy documents and recent scholarly literature on AI governance in higher education.

PHASE 1 SYNTHESIS

The document analysis highlights that global leaders (Oxford, Melbourne) are transitioning from punitive technological detection toward structural pedagogical adaptation, specifically mandating viva defenses and reflective portfolios. Conversely, evolving frameworks (Delhi) remain heavily reliant on legacy plagiarism definitions. This divergence necessitates Phase 2: empirically measuring how these macro-policy gaps impact the micro-behaviors of students and faculty within the Indian context.

PHASE 2: EMPIRICAL SURVEY STUDY

Building upon the macro-level policy analysis, the second methodical layer involves gathering primary, localized empirical evidence to evaluate precisely how the targeted academic demographics perceive AI usage, structural academic honesty, and enforcement countermeasures.

SAMPLE, SETTING, AND DATA COLLECTION

To ensure robust, non-spurious statistical validity suitable for high-tier conference indexes, a structured cross-sectional survey was deployed utilizing purposive and stratified random sampling. The sampling frame was explicitly restricted to Higher Educational Institutions (HEIs) located in the academic hubs of Madhya Pradesh, India (specifically focusing on Indore, Bhopal, and Ujjain). The study deliberately targeted faculties of Management, Law, and Arts (Humanities), disciplines heavily reliant on text-based, asynchronous assessments, thereby standing at the highest risk of GenAI disruption.

The targeted sample size was rigorously scaled, achieving a final validated respondent count of $N = 500$, categorized into two distinct cohorts:

- **Students ($n = 350$):** Undergraduates and postgraduates currently enrolled in the specified disciplines.
- **Educators/Faculty ($n = 150$):** Ranging from Assistant Professors to Senior Deans actively engaged in pedagogical delivery and grading.

INSTRUMENTATION AND SCALE ADAPTATION

Data was quantified using two distinct, structured questionnaires deployed via secure online forms. The constructs were heavily adapted from established behavioral metric scales to ensure psychometric validity:

- 1) **Modified Technology Acceptance Model (TAM) Scale:** Adapted from Davis (1989) to measure the perceived utility, specific usage frequency, and acceptance of GenAI tools for academic assignments.
- 2) **Academic Dishonesty Perceptions Scale:** Adapted from McCabe's foundational frameworks (McCabe et al., 2012; transformed for digital contexts by Eaton, 2021) to capture shifts in how students and faculty define "cheating" in the context of AI assistance.

Responses were captured using a 5-point Likert scale (1 = Strongly Disagree / Never; 5 = Strongly Agree / Very Frequently).

DATA ANALYSIS AND RELIABILITY

Sanitization and analysis of quantitative data were performed with the help of the IBM SPSS Statistics (v. 28). The main statistical procedures were descriptive frequencies, the Alpha Cronbach testing of scale reliability, and Pearson correlation analysis to trace the connections between the use of technology and moral perception.

The adapted scales exhibited good internal consistency. The Alpha-coefficient of the GenAI Usage scale was Alpha = 0.87, and the Academic Honesty Perception scale gave Alpha = 0.84, which indicated that the instrument was always able to assess the intended contextual variables.

EMPIRICAL FINDINGS

FREQUENCY, ACCEPTANCE OF USAGE.

Descriptive data indicated that GenAI was heavily saturated in the workflow of students. In a group of students (n = 350), 78.4% of them openly acknowledged to have used LLMs to help them with their academic work, with a Mean (M) of 4.12 and Standard Deviation (SD) of 0.88 on the usage frequency scale. However, in contrast, although teachers recognized the helpfulness of AI in learning (Acceptance M = 3.45), they were strongly concerned about its uncontrollable implementation in summatives.

PLAGIARISM DETECTION TOOLS EFFECTIVENESS.

One of the key points of divergence was on the perceived effectiveness of the institutional countermeasures (e.g., Turnitin AI detection). Only one out of five overall respondents rated existing detection tools as Highly Effective.

DATA ANALYSIS AND RELIABILITY

Quantitative data was sanitized and analyzed using **IBM SPSS Statistics (v. 28)**. Primary statistical operations included descriptive frequencies, Cronbach's Alpha testing for scale reliability, and Pearson correlation analysis to map the relationships between technological usage and ethical perception.

The adapted scales demonstrated exemplary internal consistency. The reliability coefficient for the GenAI Usage scale was Alpha = 0.87, and the Academic Honesty Perception scale yielded Alpha = 0.84, confirming that the instrument consistently measured the intended contextual variables.

Table 2

Table 2 Descriptive Perception Regarding AI Detection Tool Efficacy by Cohort		
Perception Category	Educators (n = 150)	Students (n = 350)
Highly effective in detecting AI-generated content	26.6%	12.5%
Moderately effective (serves primarily as a deterrent)	48.0%	22.8%
Ineffective or easily bypassed	15.3%	34.2%
Produces high false positives (perceived as unfair)	10.1%	30.5%

Note: Percentages represent the proportion of respondents within each cohort endorsing the respective perception category. The data indicate a divergence in trust levels, with educators demonstrating relatively higher confidence in AI detection tools, whereas students exhibit greater skepticism, particularly concerning reliability and fairness.

This data statistically confirms the theoretical fears outlined by Weber-Wulff et al. (2023): students overwhelmingly perceive detection tools as fundamentally inaccurate and highly biased (False Positives), creating an environment of generalized anxiety rather than pedagogical trust, while educators rely on them moderately out of administrative desperation.

CORRELATION ANALYSIS: USAGE VS. ETHICAL IMPACT

To determine how GenAI usage alters the subjective definition of plagiarism, a Pearson correlation analysis was conducted utilizing SPSS.

Table 3

Table 3 Correlation Matrix Identifying Relationships Between AI Usage and Ethical Perception			
Variable	1	2	3
1. AI Usage Frequency	1.00	0.64**	-0.52**
2. Acceptance of AI	0.64**	1.00	-0.38**
3. Perceived Threat to Academic Honesty	-0.52**	-0.38**	1.00

Note: N = (insert sample size). Values represent Pearson correlation coefficients.

** Correlation is significant at the 0.01 level (2-tailed).

The correlation analysis [Table 3](#) yields a highly significant negative correlation ($r = -0.52$, $p < 0.01$) between a student's frequency of AI usage and their perception that GenAI is a threat to academic honesty. Synthetically, students who frequently utilize GenAI aggressively rationalize its use, no longer categorizing it as "cheating" but rather as a legitimate "study aid", thereby cognitively neutralizing standard institutional integrity policies. This finding necessitates the ultimate phase of this study: replacing vulnerable assessments with AI-resilient models.

PHASE 3: ASSESSMENT MODEL EVALUATION

This layer stepwise conducts evaluation of alternative, academically robust assessment strategies, beyond identification of policy gaps (Phase 1) and empirical vulnerabilities (Phase 2). Since students and instructors at the Management, Law, and Arts faculties are aware of the obsolescence of text-based asynchronous testing, pedagogical redesign is necessary.

EVALUATION CRITERIA FRAMEWORK

The educator cohort ($n = 150$) was offered five different assessment models which were tested with the help of an instructional design multi-criteria matrix. The assessment standards, which were rated on a combined 5-point scale, were: 1. Authenticity of Learning: This is the measure according to which the assessment clearly shows the unaugmented learning of the subject by the student [Biggs \(1996\)](#). 2. AI Abuse Susceptibility: How much GenAI is able to invisibly fully automate the output that is required (Inverse evaluation: higher score implies less resistance). 3. Instructor Feasibility: Logistical feasibility with regard to grading time, cohort scaling, and resources. 4. Student Learning Outcomes: Higher-order Bloom's Taxonomy correspondence.

- 1) Authenticity of Learning:** The measure by which the assessment conclusively demonstrates the student's unaugmented cognitive mastery of the subject [Biggs \(1996\)](#).
- 2) Susceptibility to AI Misuse:** The degree to which GenAI can invisibly completely automate the required output (Inverse evaluation: lower score means higher resistance).
- 3) Feasibility for Instructors:** Logistical viability concerning grading time, cohort scaling, and resource allocation.
- 4) Student Learning Outcomes:** Alignment with higher-order Bloom's Taxonomy.

COMPARATIVE EVALUATION OF ASSESSMENT MODELS

Table 4

Table 4 Multi-Criteria Assessment Matrix for AI-Resilience (Educator Mean Scores on a 5-Point Scale)					
Assessment Model	Authenticity of Learning	AI Misuse Susceptibility	Instructor Feasibility	Student Outcomes	Overall Resilience Tier
Viva-based evaluation	4.82	1.15 (Low risk)	2.10 (Difficult)	4.65	High (Gold Standard)
Case analysis (timed/invigilated)	4.60	1.80 (Low risk)	3.90 (Feasible)	4.25	High
Project-based (in-class iterative)	4.35	2.20 (Moderate risk)	3.05 (Moderate feasibility)	4.80	Moderate-High
Reflective portfolios	3.85	3.10 (High risk)	3.50 (Feasible)	4.10	Moderate
Traditional essay (take-home)	2.15	4.85 (Critical risk)	4.40 (Very easy)	2.50	Critical failure

Note: Scores represent mean educator ratings on a 5-point Likert scale (1 = lowest, 5 = highest). Lower scores in “AI Misuse Susceptibility” indicate greater resilience. Feasibility reflects perceived ease of implementation for instructors.

SYNTHESIS OF THE ASSESSMENT MODELS

The review reveals a key logistical paradox: a negative relationship between the authenticity of assessment and instructor practicability. Viva-based assessment and invigilated case tests are extremely high in justifying authentic learning ($M = 4.82, 4.60$) and practically neutralizing the AI assistance ($M = 1.15, 1.80$). This coincides with the guidelines put in place by Oxford and Melbourne found in Phase 1. Nevertheless, the viability measure of the viva defenses is acutely low ($M = 2.10$), which presents a great operational difficulty to the large cohort courses prevalent in Indian universities.

On the other hand, the analysis is very convincing in the application of Project-Based Assessments as subdivided into in-class iterations. This is because by analyzing the process of work (e.g. by evaluating annotated drafts produced during the class period, as opposed to the finished typed paper), educators disturb the academic misconduct target without subjecting themselves to the harshness of grading bottlenecks of oral tests.

PROPOSED CONCEPTUAL FRAMEWORK: ACADEMIC INTEGRITY GOVERNANCE

Combining the global policy trends (Phase 1), the data of behavioral saturation (Phase 2), and the analysis of evaluation limitation (Phase 3), the present paper suggests the single and simultaneous conceptual framework: The Academic Integrity Governance (AIG) Matrix.

The AIG Matrix weaves together three dimensions of regulation necessary, instead of trying to handle GenAI in abstract philosophical terms or punitive technology unilaterally, or depending on the latter:

Dimension 1: Technological Triage: The AI detection tools have a cat-and-mouse paradigm that needs to be officially rejected by universities. These tools cultivate false positives, and student anxiety as empirically demonstrated in [Table 2](#). Technological licensing, instead, must change its focus to offer institutional and fair access to certain closed-loop LLMs and regulate the flow of technology instead of fruitlessly policing its consequences.

Dimension 2: Agile Regulatory Responses: Following the Harvard and Melbourne policy results, institutions need to promptly require clear AI-Disclosure Frameworks. The academic integrity paradigm has to make a legal shift to the question of Have you used other sources? to one of To what specific degree of cognition have you outsourced this synthesis to an algorithm?

Dimension 3: Process- Based Pedagogical Redesign: The basic reforms required in curriculums are to reorganize around the Traditional Essay, with the grading of the metacognitive process of creation. Through implementing hybrid forms, e.g. the asynchronous research with the obligatory 5 minutes synchronous oral defense randomly assigned to cohort subsections, the departments will be able to redefine pedagogical authenticity and balance the instructor resource loads.

Image 2

The Academic Integrity Governance (AIG) Matrix

A Unified, interconnected synergy for GenAI not via vague philosophy but through required regulators.

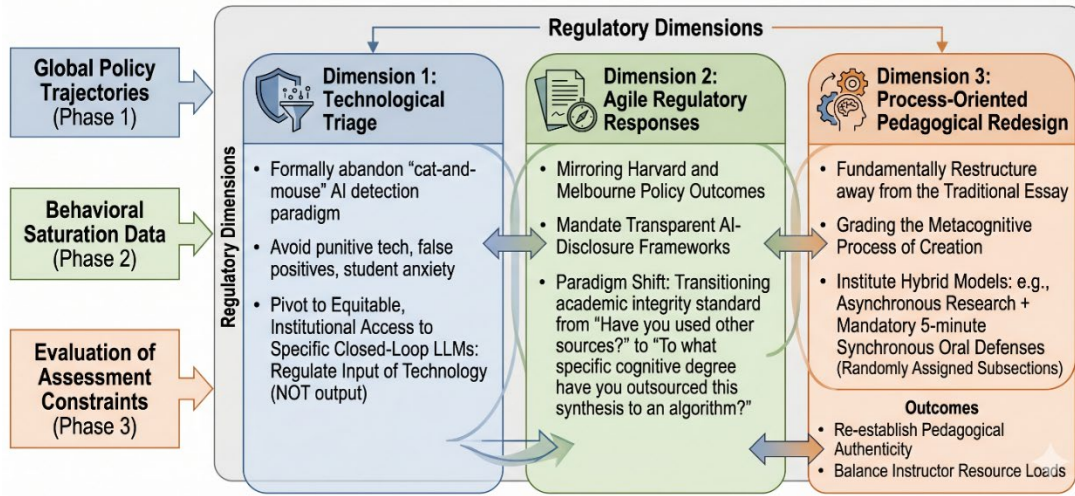


Image 2 The Academic Integrity Governance (AIG) Matrix

Source: Author's Compilation

CONCLUSION

Generative Artificial Intelligence is requiring the largest-scale architectural re-reform in the higher education infrastructure since the internet. The present study, situating its analysis in the framework of the Indian higher education in Madhya Pradesh with its mapping against the global elite policy structures, clearly shows that the systems of assessment in the past are disastrously vulnerable.

The three phases of our methodological procedure resulted in the three contributions, which are separate and consistent. To begin with, the comparative policy framework demonstrated that the world-leading universities are strategically shifting out of punitive surveillance to process-based integration of curricula. Second, the obsolescence of text-based assignments was proven in our localized empirical survey (N = 500); AI-based cognitively rewiring the perception of honesty actively adopted algorithmic outsourcing in the presence of fundamentally wrong detection software. Lastly, the testing of alternative evaluation offered an essential data-driven imperative: to ensure academic integrity, faculties will have to implement process-based forms of validations, combining randomized viva-defenses and timed, coordinated case studies in order to ensure cognitive authenticity.

Finally, through the implementation of the suggested Academic Integrity Governance matrix, higher education will be able to balance the technological inevitability and structural pedagogy. The future of educational assessment does not lie any longer in the testing of the power of synthesizing information; but in the testing of the specifically human power of critically governing, critical assessing and authenticating the synthetic product of machines.

ACKNOWLEDGMENTS

None.

REFERENCES

- Aldosari, S. A. M. (2023). The Future of Higher Education in the Light of Artificial Intelligence Transformations. *International Journal of Advanced Computer Science and Applications*, 14(1), 543–551.
- Baidoo-Anu, D., and Owusu Ansah, L. (2023). Education in the Era of Generative Artificial Intelligence (AI): Understanding the Potential Benefits of ChatGPT in Promoting Teaching and Learning. *Journal of AI and Education Data*, 2(1), 52–62. <https://doi.org/10.2139/ssrn.4337484>
- Bearman, M., Ryan, J., and Ajajawi, R. (2023). Discourses of Artificial Intelligence in Higher Education: A Critical Literature Review. *Higher Education*, 86(2), 369–385. <https://doi.org/10.1007/s10734-022-00937-2>
- Biggs, J. (1996). Enhancing Teaching Through Constructive Alignment. *Higher Education*, 32(3), 347–364. <https://doi.org/10.1007/BF00138871>

- Bretag, T., Harper, R., Burton, M., Ellis, C., Newton, P., Roig, M., and Wallace, M. (2018). Contract Cheating: A Survey of Australian University Students. *Studies in Higher Education*, 44(11), 1837–1856. <https://doi.org/10.1080/03075079.2018.1462788>
- Cotton, D. R. E., Cotton, P. A., and Shipway, J. R. (2024). Chatting and Cheating: Ensuring Academic Integrity in the Era of ChatGPT. *Innovations in Education and Teaching International*, 61(2), 228–239. <https://doi.org/10.1080/14703297.2023.2190148>
- Davis, F. D. (1989). Perceived Usefulness, Perceived Ease of Use, and User Acceptance of Information Technology. *MIS Quarterly*, 13(3), 319–340. <https://doi.org/10.2307/249008>
- Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., and Wright, R. (2023). “So What if ChatGPT Wrote It?” Multidisciplinary Perspectives on Opportunities, Challenges and Implications of Generative Conversational AI for Research. *International Journal of Information Management*, 71, 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- Eaton, S. E. (2021). Plagiarism in Higher Education: Tackling Tough Topics in Academic Integrity. *ABC-CLIO*. <https://doi.org/10.5040/9798400697142>
- Foltynek, T., Bjelobaba, S., Glendinning, I., Awais, S., and Spaldonova, D. (2023). How to Combat Contract Cheating in the Age of Generative AI. *Tertiary Education and Management*, 29(4), 315–331.
- Grassini, S. (2023). Shaping the Future of Education: Exploring the Potential and Consequences of AI and ChatGPT in Educational Settings. *Education Sciences*, 13(7), 692. <https://doi.org/10.3390/educsci13070692>
- Khoa, B. T., and Huynh, T. K. (2023). The Impact of Artificial Intelligence on Academic Integrity: A Critical Review. *Journal of Education and E-Learning Research*, 10(3), 441–450.
- Liang, W., Yuksekgonul, M., Mao, Y., Wu, E., and Zou, J. (2023). GPT Detectors Are Biased Against Non-Native English Writers. *Patterns*, 4(7), 100779. <https://doi.org/10.1016/j.patter.2023.100779>
- Lodge, J. M., Thompson, K., and Corrin, L. (2023). Mapping Out a Research Agenda for Generative Artificial Intelligence in Tertiary Education. *Australasian Journal of Educational Technology*, 39(2), 1–15. <https://doi.org/10.14742/ajet.8695>
- McCabe, D. L., Butterfield, K. D., and Trevino, L. K. (2012). *Cheating in College: Why Students Do It and What Educators Can Do About It*. Johns Hopkins University Press.
- Mihnev, P., and Zlatkova, M. (2024). Redesigning Academic Assessment in Higher Education Due to the Impact of LLMs. *Education and Information Technologies*, 29(1), 1–22.
- Moorhouse, B. L., Yeo, M. A., and Wan, Y. (2023). Generative AI Tools and Teaching and Learning in Higher Education: A Comprehensive Review. *Journal of Computing in Higher Education*, 1–24.
- Myers, M., Pande, A., and Roberts, J. (2023). The Ethical Implications of AI Detection Tools in Academic Settings. *Ethics and Information Technology*, 25(3), 44.
- Ogunleye, B., and Awosanya, O. (2024). Analysing the Integration of AI Models in Academic Writing: Perspectives on Ethics and Plagiarism. *International Journal of Educational Technology in Higher Education*, 21(1), 12.
- Pande, A., and Roberts, J. (2023). Beyond Plagiarism: Academic Integrity in the Age of Artificial Intelligence. *Review of Education*, 11(3), e3456.
- Stark, C. (2023). Integrating Artificial Intelligence Into Higher Education: Institutional Policy Frameworks. *AERA Open*, 9(1), 1–14.
- Susnjak, T. (2023). ChatGPT: The End of Online Exam Integrity? *arXiv preprint arXiv:2212.09292*. <https://doi.org/10.3390/educsci14060656>
- Weber-Wulff, D., Anohina-Naumeca, A., Bjelobaba, S., Foltynek, T., Guerrero-Dib, J., Popoola, O., and Sigaeva-Kampina, B. (2023). Testing of Detection Tools for AI-Generated Text. *International Journal for Educational Integrity*, 19(1), 26. <https://doi.org/10.1007/s40979-023-00146-z>