



The Impact of Artificial Intelligence on Academic Integrity: A Scoping Review

Alexa Clarisse N. Loberes¹, ORCID No. 0009-0003-3538-5297

John Paulo G. Doroteo², ORCID No. 0009-0001-6252-3603

¹Discipline Education Officer, De La Salle University, 2401 Taft Avenue, Manila, Philippines

²STREAM & Robotics Faculty, De La Salle University Integrated School, BiEan City, Laguna, Philippines

Article History:

Initial submission:	20 October 2025
First decision:	23 October 2025
Revision received:	20 December 2025
Accepted for publication:	23 December 2025
Online release:	31 December 2025

Abstract

In the post-pandemic era, the rise of Generative Artificial Intelligence (GenAI) has sparked significant concerns from educators and researchers regarding academic integrity. These concerns are highlighted by a 2024 study by Waltzer et al., which illustrates the growing tension between technological innovation and ethical standards in university classrooms. This scoping review systematically maps existing literature to identify how student use of GenAI aligns with documented integrity risks and institutional responses. Following PRISMA-ScR standards, a total of 281 articles were retrieved from Scopus and EBSCOhost between 2023 and 2025, resulting in a finalized dataset of 35 empirical studies. Data extraction utilized the SAMR (Substitution, Augmentation, Modification, Redefinition) model as a diagnostic framework to measure the depth of technology integration. Results indicate a significant demographic concentration, with 88.6% of research situated in Higher Education, and ChatGPT identified as the primary tool utilized. Analysis reveals that while 68.6% of usage falls within the Enhancement phase, there is a critical mismatch between "Process-oriented" risks, such as cognitive erosion, and current "Product-oriented" institutional solutions like detection tools. To address this gap, the study proposes the original SAMR-Integrity-Response (SIR) Framework. This model provides a strategic roadmap for educational institutions, advocating for a shift from defensive regulatory postures at lower integration levels to evolutionary pedagogical pivots, including assessment redesign and process-based grading, at transformative levels. This review equips institutions with the tools to preserve integrity in digitally enabled learning environments.

Keywords: Generative Artificial Intelligence, LLM, Large Language Model, Academic Integrity, Academic Dishonesty



Copyright © 2025. The Author/s. Published by VMC Analytikis Multidisciplinary Journal News Publishing Services. The Impact of Artificial Intelligence on Academic Integrity: A Scoping Review © 2025 by Alexa Clarisse N. Loberes and John Paulo G. Doroteo is an open access article licensed under **Creative Commons Attribution (CC BY 4.0)**. This permits the copying, redistribution, remixing, transforming, and building upon the material in any medium or format for any purpose, even commercially, provided that appropriate credit is given to the copyright owner/s through proper and standard citation.

INTRODUCTION

The COVID-19 pandemic served as a global catalyst, compressing a decade's worth of digital evolution into a brief period and fundamentally resetting the technological baseline of modern classrooms (UNESCO, 2023). This established a digital ecosystem characterized by high data availability and ubiquitous connectivity, which has paved the way for the current rise of Artificial Intelligence, shifting the focus from simple digital access to the complex application of AI-driven personalized learning and automated instruction. As institutions achieved a newfound 'digital maturity,' AI emerged not as a separate tool, but as the natural next frontier for

personalizing instruction and automating academic management in the post-pandemic era (Arias Ortiz et al., 2025). Building upon this rapid digitalization, Wang et al. (2024) report that Gen AI technology and educational robotics are now deemed integral to the learning and training management systems, which aid and automate multiple processes in the educational system and activities of teaching and learning.

Despite the changes in the educational landscape, the commitment to delivering quality education remains constant. Quality education is predicated on academic integrity, an overarching framework encompassing the foundational principles of honesty, trust, and ethical behavior (Balalle & Pannilage, 2025).

While academic integrity is a long-standing pillar in the academe, the shift to digital environments has introduced new vulnerabilities. Research indicates that online assessments are associated with an increased frequency of academic dishonesty cases (Roe et al., 2023). This risk is further intensified by GenAI, which presents a new challenge to traditional integrity standards (Balalle & Pannilage, 2025). GenAI utilizes machine learning models trained on vast datasets to identify underlying patterns, enabling the production of original text, visual, or auditory content based on user inputs. This process differs fundamentally from traditional search engines; rather than retrieving and organizing existing information through indexing, these tools synthesize entirely new outputs by predicting the next logical element in a sequence (UCTL, n.d.).

In the current educational setting, GenAI significantly escalates this academic integrity crisis. The ability to generate content without human input raises complex ethical questions regarding the line between human intellectual effort and machine substitution. Furthermore, GenAI challenges traditional definitions of intellectual work as issues like authorship and ethical co-creation broaden the scope of misconduct beyond simple plagiarism. (Ellis et al., 2023) These tools facilitate sophisticated forms of deception, such as 'illicit paraphrasing' to evade detection and the 'AI ghostwriter effect,' where students misrepresent machine-generated content as their own work (Draxler et al., 2023). Ultimately, the severity of these integrity breaches is directly linked to the degree to which the technology has replaced the student's required intellectual work. As educational institutions adopt AI technologies, it is crucial that they implement structural changes to uphold ethical standards.

Existing reviews often provide a superficial view of AI involvement, failing to adequately map the extent of integration against the erosion of intellectual work. To address this gap, this study utilizes the SAMR model, as it provides a fundamental justification for

measuring the depth of GenAI integration—a factor the researchers argue is directly proportional to the erosion of academic integrity.

The aim of this scoping review is to systematically map the existing literature on academic integrity within the context of post-pandemic adoption of GenAI. Data will be mapped using the SAMR (Substitution, Augmentation, Modification, Redefinition) model. Created by Puentedura, the SAMR model is a four-level taxonomy that serves as a means of assessing and integrating technology in education. As technology is primarily integrated into education to enhance and improve teaching and learning, the model encourages an upward movement through the levels of teaching technology (Blundell et al., 2022). By employing this ladder taxonomy, the review advances beyond simple detection to distinguish between minimal-impact substitutions and maximal-impact transformations, establishing a prescriptive model for institutional intervention.

This study aims to provide an integrative examination of the literature, encompassing various research streams, to articulate core ideas, approaches, and areas of underdeveloped research that are relevant to educators, educational institutions, and educational technologists interested in fostering a culture of academic integrity in digitally enabled learning environments.

Objectives. The objectives below outline the study's focus on generative AI use, institutional responses, and integrity safeguards:

1. To systematically identify and categorize the ways in which the literature reports the uses of generative AI among high school and college students and determine the frequency distribution across SAMR levels.
2. To specify types of institutional responses concerning the use of generative AI.
3. To propose an original prescriptive model that maps integrity risks to institutional

interventions across the SAMR levels, establishing a novel framework for mitigating AI-driven academic misconduct.

LITERATURE REVIEW

Academic Integrity and Misconduct in the Digital World. Academic institutions uphold the principles of academic integrity because the goals of teaching, learning, research, and service can only be achieved in an environment that fosters ethical conduct. However, there is often a disconnect between this ideal and institutional practice. The International Center for Academic Integrity (ICAI) notes that scholarly organizations frequently struggle to describe their commitment to integrity in positive, practical terms, often focusing instead on reactive policies that prohibit specific types of misconduct. To foster a more proactive culture, the ICAI offers a framework of six fundamental values: honesty, trust, fairness, respect, responsibility, and courage (ICAI, 2021, p. 4). These values are more than abstract principles; they are intended to enable academic communities to translate ethical ideals into action and improve decision-making behavior.

In the current landscape of Generative AI, these values serve as essential guideposts for navigating new ethical challenges. *Honesty & Trust* requires being truthful and free from fraud or deception. In a GenAI context, this is undermined when students misrepresent machine-generated content as their own genuine work, thereby breaking the assured reliance on the truth of student capability that the academic community requires to function. *Fairness* involves impartial treatment and the expectation that all members do their own original work. The use of GenAI to gain an unmerited advantage violates this principle. Furthermore, *Respect* is demonstrated through the "proper identification and citation of sources," a standard that is often complicated by GenAI's ability to obscure intellectual contributions. *Responsibility* demands that individuals be accountable for their own actions and follow institutional rules even when peer

pressure or new technology makes misconduct easier. Finally, *Courage* is the capacity to act on these values despite fear. For students today, this means choosing to maintain integrity even when facing the risk of negative consequences, such as a lower grade, compared to peers using AI tools.

Although these values remain constant, their application must be context-sensitive (Bretag et al., 2019). In the case of GenAI, context sensitivity requires acknowledging that the digital environment has fundamentally changed the student-instructor relationship. The transition to remote learning during the COVID-19 pandemic exposed significant vulnerabilities in conventional assessment models, particularly during online evaluations (Ratten, 2023; Roe, et. al 2023). This created an imbalance in the academic ecosystem, where the accessibility of generative tools for academic dishonesty far outpaced the institutional capacity for proctoring, oversight, and pedagogical adaptation.

GenAI tools, particularly Large Language Models (LLMs), have escalated the academic integrity crisis (Balalle & Pannilage, 2025). The definition of intellectual work has become increasingly complex as the traditional binary of "original vs. plagiarized" fails to capture modern nuances. Rather than simply broadening the scope of existing work, GenAI introduces entirely new categories of academic misconduct, such as "ghost-authoring" through iterative prompting and the obfuscation of ethical co-creation (Morris et al., 2023). This ability to generate high-fidelity content without primary human input forces organizations to re-evaluate the boundary between human intellectual effort and machine substitution (Ellis et al., 2023).

A framework is required that moves beyond policing superficial AI usage to evaluating the depth of machine integration in the creative process.

The SAMR Framework. Addressing the challenges posed by GenAI requires a

framework to evaluate the extent to which GenAI undermines student intellectual engagement. This assessment is based on the premise that the risk to academic integrity is directly linked to the degree of machine substitution, where higher levels of automated content generation correspond to a significant decrease in authentic human effort. This relationship provides a fundamental justification for integrating current literature through the lens of the SAMR (Substitution, Augmentation, Modification, Redefinition) model, allowing for a nuanced categorization of GenAI's role in the creative process.

The SAMR model—comprising of Substitution, Augmentation, Modification, and Redefinition—was originally developed to support educators in selecting technologies that enhance the learner experience and facilitate higher-order thinking (Rehman & Aurangzeb, 2021). While traditionally used for instructional design, the model is uniquely applicable to academic integrity analysis because it provides a taxonomy for the "depth of integration" between the student and the machine. By mapping GenAI usage onto these four levels, researchers can identify the threshold at which AI ceases to be a supportive tool (Augmentation) and begins to substitute for the student's cognitive agency (Modification/Redefinition). In this context, SAMR serves as a diagnostic framework to evaluate the "displacement of authorship"—the point at which the technological integration becomes so deep that it effectively erases the student's intellectual contribution. SAMR, developed by Puentedura (2006), is divided into two layers, each with two levels. The Substitution and Augmentation Levels under the Enhancement Layer, and the Modification and Redefinition Levels are under the Transformational Layer.

Enhancement:

Substitution is where technology is used to solely replace or replicate a previous pedagogical practice without any change to the function or purpose of that practice. Ethical substitution includes using tools for rapid

information retrieval to scan background literature or for basic language translation (Gruenhagen et al., 2024). However, this level presents high-risk integrity issues when the tool is used for AI-based/assisted plagiarism or contract cheating, where the technology entirely substitutes the student's role (Alawad et al., 2025).

Augmentation focuses on improving access to educational information; however, the same instructional practice is still used. Ethically, this includes using GenAI for brainstorming, proofreading, or outlining difficult texts (Caling et al., 2025). The corresponding integrity risk involves intentional concealment and illicit paraphrasing, where students use GenAI's sophisticated rewriting capabilities to evade automated content detectors (Waltzer et al., 2024)

Transformation:

In **Modification**, a change in the learning experience occurs through the integration of technology and a significant redesign of the instructional task. For example, GenAI acts as an interactive tutor or "co-pilot," providing immediate, personalized feedback that allows students to "see logical gaps" in their arguments (Akiba & Garte, 2024)

In **Redefinition**, the highest level of integration, technology is utilized to accomplish tasks that were previously impossible. In the GenAI landscape, this is realized through Personalized Learning, where the technology adapts content precisely to an individual student's unique learning journey (Dwivedi et al., 2023; Kofinas et al., 2025).

The SAMR model provides a structural foundation for analyzing how GenAI is integrated into academic work. By distinguishing between minimal-impact enhancements and maximal-impact transformations, the model clarifies how the role of the student shifts from "primary author" to "prompt engineer." More significantly, this framework highlights the threshold of cognitive

delegation, marking the moment when a task’s transformation no longer depends on the student’s growing mastery of a tool. Instead, it emerges from the machine’s capacity to substitute intellectual effort, reshaping learning into mediated automation.

METHODS

This report outlines a scoping review that has systematically mapped the literature on the impact of Generative Artificial Intelligence (AI) technologies on academic integrity. The review adhered to the PRISMA-ScR (Preferred Reporting Items for Systematic Reviews and Meta-Analyses for Scoping Reviews) standards for design and reporting, ensuring transparency and reproducibility. To ensure a focused research scope, the objectives and inclusion/exclusion criteria were developed using the PCC (Population, Concept, Context) framework. This framework guides the review by identifying the specific groups being studied (Population), the core phenomenon of GenAI and academic integrity (Concept), and the academic or institutional settings involved (Context).

The inclusion and exclusion criteria were established based on the PCC framework to define the scope of the search:

Table 1
Inclusion-Exclusion Criteria using PCC Framework

Filter Category	Criterion	Inclusion (KEEP)	Exclusion (DISCARD)
1. Population	P-Filter	Focuses on High School and/or College/University students.	Focuses solely on faculty, staff, administrators, K-8 students, or industry/corporate settings.
2. Concept	C-Filter (Technology)	Explicitly discusses Generative AI (GenAI) or LLMs.	Discusses only general technology (e.g., standard plagiarism checkers, pre-2023 AI)
3. Context	Co-Filter (Topic)	Discuss academic integrity by describing or focusing on a specific task or way students use GenAI, or focusing on a specific intervention, policy, or assessment change.	Focused on establishing frameworks, vague commentaries, speculative essays, editorials, opinion pieces, letters, or news articles. Discuss the existence of AI but provide no specific evidence of actual student behavior or specific institutional solutions.

Information Sources. Data retrieval for this scoping review commenced in March 2025 and was finalized prior to the submission of this manuscript in November 2025. This study

utilized two primary platforms: Scopus and EBSCOhost. Scopus was selected for its status as a premier multidisciplinary database with extensive coverage of the social sciences, providing a high-level overview of the global research landscape. To ensure a comprehensive search of education-specific literature, EBSCOhost was used to simultaneously query specialized indexes, including ERIC (Education Resources Information Center) and Education Source. This combination ensured that both broad multidisciplinary perspectives and deep pedagogical research on the ethics of education were captured.

Search Strategy. The search strategy was designed to be comprehensive, ensuring the inclusion of diverse literature across the three core concepts of this study: the student population, generative artificial intelligence, and academic integrity. Following the building block approach, keywords were categorized into three concepts, and synonyms were linked using the Boolean operator “OR” to maximize the retrieval of relevant articles. These concepts were then combined using the “AND” operator to focus the search on the intersection of these topics.

To ensure the review reflects the most current landscape of Generative AI in education, the search was limited to articles published between 2023 and 2025. This timeframe captures the rapid development and implementation of AI tools following the public release of major large language models. No language restrictions were initially applied, although the final selection was limited to English-language publications to ensure accurate thematic synthesis.

To ensure the sensitivity and specificity of the search strategy, a pilot test was conducted in EBSCOhost and Scopus. The initial search string was tested to evaluate the relevance of the first 20 hits against the study’s inclusion criteria. Initial results were screened to ensure they captured the intersection of student behavior and AI-related integrity risks. The search

strategy was finalized once the pilot results yielded a high proportion of relevant studies, ensuring that the search was sufficiently broad to capture diverse literature while maintaining focus on the research objectives.

In addition to the electronic database search, the researchers performed a manual 'hand search' of the reference lists of all final included studies. This backward snowballing technique was used to ensure that any key foundational or empirical studies not captured by the initial Boolean search were identified and screened for inclusion. A total of 281 articles were exported, organized, and tagged in Google Sheets to undergo the title screening process.

Title Screening Process. The initial records were exported to Google Sheets for organization and a preliminary title screening to ensure alignment with the research objectives. Following this aggregation, a two-step deduplication process was conducted: 15 duplicates were manually removed to streamline the dataset, followed by a secondary verification using the 2025 version of Rayyan AI. Rayyan AI is a specialized, web-based systematic review tool that ensures the accuracy of the final dataset and provides a transparent audit trail for the autonomous review phase, thereby adhering to the transparency and reliability standards of the PRISMA-ScR protocol. The researchers utilized Rayyan AI only during deduplication process.

Abstract Screening Process. The remaining 266 records underwent a manual abstract screening process. The researchers independently screened the abstracts, applying the pre-defined inclusion and exclusion criteria. Decisions were recorded in a structured Google Sheets matrix for each respective researcher. This log included the final decision and the specific rationale for exclusion based on PCC.

Initial independent screening resulted in a consensus on 222 records (174 joint inclusions and 48 joint exclusions). Disagreements occurred in 44 cases, where 39 studies were included by Researcher A but excluded by

Researcher B, and 5 studies were included by Researcher B but excluded by Researcher A.

These results yielded an initial Cohen's Kappa (κ) of 0.5540, representing moderate agreement according to the benchmarks established by Landis and Koch (1977) [1.1]. To ensure the highest level of methodological rigor, all 44 discordant records were moved to a formal reconciliation phase. During this phase, the two researchers met to discuss each discordant study in detail, evaluating its relevance to specific GenAI use or misuse, as well as potential institutional interventions. This collaborative process ensured 100% consensus on the final list of studies that would move forward to the full-text eligibility assessment. After the reconciliation phase, a total of 79 studies were excluded, and 187 studies were included based on the established criteria.

Full Text Retrieval and Review. Following the abstract screening, 187 studies were sought for full-text retrieval. Of these, 43 reports could not be retrieved due to a lack of institutional access, inactive links, or being behind restrictive paywalls despite attempts to locate open-access versions.

Consequently, 144 studies underwent a rigorous full-text eligibility assessment. During this stage, each paper was read in its entirety to ensure it provided the necessary data to address the research objectives, specifically the presence of student behavioral descriptions for SAMR mapping and documented institutional interventions. The full-text articles were screened using a collaborative consensus-based approach involving both researchers synchronously. Rather than independent screening followed by a tie-breaker, the researchers conducted a joint review of each study to ensure maximum semantic alignment and consistency in applying the inclusion criteria.

First, each study was screened by the researchers for its structural evidence base. The researchers prioritized primary research and case studies over conceptual or speculative

literature. To be retained, a study had to meet three benchmarks: an explicit identification of high school or higher education student participants, a documented research design (e.g., qualitative case study, action research, or empirical observation), and a clear description of the specific GenAI tool used and the nature of the student task. 58 studies were removed as they did not meet the benchmarks that were set, and they would often lack empirical student data.

Next, the studies that passed the structural filter were then subjected to a detailed data assessment. This stage was designed to extract the evidence necessary to fulfill the study's core objectives. Each text was analyzed for descriptions of functional changes in student workflows, allowing for categorization into Substitution, Augmentation, Modification, or Redefinition. Then, researchers identified specific policy shifts, assessment redesigns, or pedagogical interventions triggered by the use of AI in this case. Of the remaining 86 studies, 26 were excluded because they lacked clear descriptions that could be categorized in the SAMR framework or did not identify specific policies, assessment redesigns, or interventions.

Following this two-stage full-text screening, 60 studies were categorized based on their utility for the Prescriptive Model. Studies providing a direct link between a specific student behavior and an institutional response were labeled as "High Priority" for the final synthesis.

Data Extraction and Charting. To ensure consistency and alignment with the study's four objectives, a data extraction matrix was developed. Consistent with the principles of evidence-based synthesis (Tricco et al., 2018), 25 studies with incomplete reporting were excluded to prevent the introduction of speculative data into the prescriptive model, thereby ensuring the framework's practical utility for educational institutions. A total of 35 studies were included in the data extraction matrix (Figure 1).

This standardized charting tool was designed to capture both descriptive study characteristics and functional data points necessary for the proposed prescriptive model. The extraction matrix was developed in accordance with the JBI Manual for Evidence Synthesis (Aromataris et al., 2024)[2.1][3.1][4.1][5.1][6.1][7.1]. To ensure transparency and reproducibility, the key requirements of the PRISMA-ScR standards (Tricco et al., 2018)[8.1][9.1][10.1] were specifically structured to capture variables that directly inform the research objectives. Following the recommendations of Levac et al. (2010)[11.1][12.1][13.1], the researchers conducted a pilot extraction on 10% of the studies to calibrate the SAMR mapping definitions and ensure inter-rater reliability. The matrix was populated through a collaborative consensus-based approach (Arksey & O'Malley, 2005)[14.1][15.1][16.1] in Google Sheets, where both researchers reviewed the full texts synchronously to resolve ambiguities regarding functional Generative AI tasks and institutional responses.

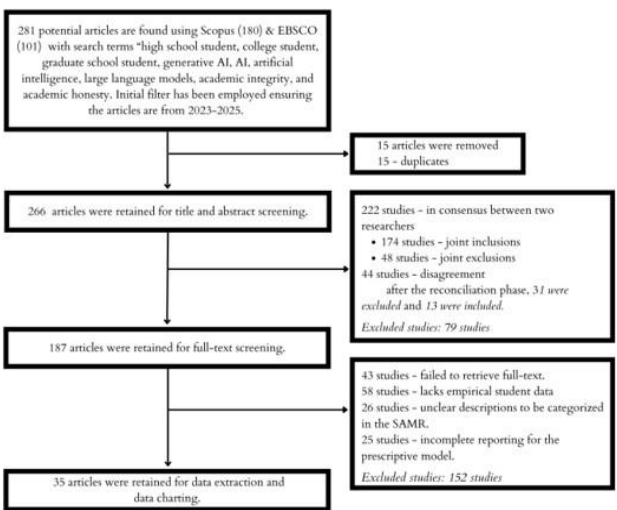


Figure 1
Search and scoping procedure

The matrix was organized into five critical data domains, each mapped to a specific research objective:

1. Extraction of the specific Generative AI tool (e.g., LLMs, image generators) and the exact nature of the student task.

2. Manual classification of the AI's role into Substitution, Augmentation, Modification, or Redefinition based on functional task changes.
3. Identification of documented misconduct types (e.g., unauthorized content generation, ghostwriting) associated with each use case.
4. Recording of specific pedagogical or policy responses implemented by the institutions.
5. Categorization of responses into Regulatory, Pedagogical, or Technical frameworks.

RESULTS

Descriptive Landscape of Generative AI in Education [17.1]. The analysis revealed a significant geographic and demographic concentration in the current literature. The data reveal a significant disparity in research focus between educational levels. As shown in Figure 2, 88.6% (n=31) of the studies focused on Higher Education (College/University), while only 11.4% (n=4) addressed High School contexts.

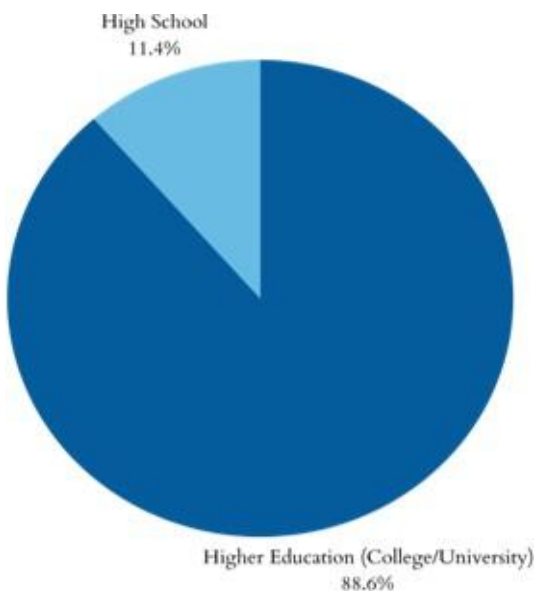


Figure 2
Educational setting

This suggests that the current empirical understanding of GenAI and academic integrity is heavily weighted toward adult learners. This demographic concentration highlights a critical

research gap, necessitating further investigation into the impact of GenAI on academic integrity within secondary education and K-12 contexts.

As shown in Figure 3, a total of 16 tools with 61 mentions were recorded across the 35 studies, indicating that many research contexts involve multiple platforms. ChatGPT is the undisputed leader in the literature, appearing in nearly every study. The presence of tools like Grammarly and Quillbot alongside LLMs suggests the interaction between traditional assistive technologies and generative agents in student workflows.

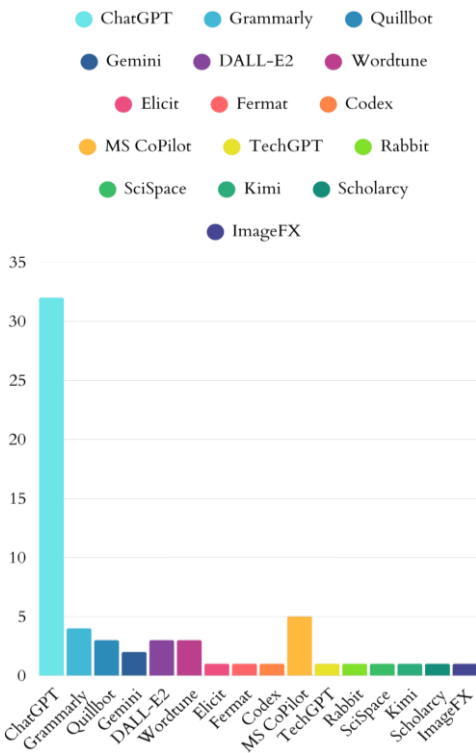


Figure 3
GenAI tools used

The reported student tasks in the data extraction matrix were synthesized using an inductive thematic analysis where themes and patterns emerge organically from the data itself, rather than being imposed by a pre-existing theory (Braun & Clarke, 2006). The researchers followed a six-step process: (1) data familiarization through matrix charting, (2) initial code generation for each specific task, (3)

searching for themes among the codes, (4) reviewing themes against the full 35-study dataset, (5) defining and naming themes, and (6) producing the final report.

This thematic analysis grouped the studies' specific AI use into three functional themes:

Theme 1 - Generative Tasks. This includes the high-level creation of "new" content, such as generating full essays, writing code/programs, drafting assignments, and developing initial ideas.

Theme 2 - Linguistic Refinement. These tasks focus on modifying existing student work. Reported actions include proofreading, revising, editing, grammar checking, paraphrasing, and summarizing texts.

Theme 3 - Cognitive Support. This category involves the "pre-writing" or research phase, including brainstorming, outlining, and seeking feedback, comments or suggestions on academic work.

The resulting three functional themes: Generative, Linguistic Refinement, and Cognitive Support, were developed based on the functional role the AI tool played within the academic workflow. This categorization was specifically chosen to facilitate the subsequent mapping to the SAMR framework, as it allows for a clear distinction between AI as a direct replacement for human effort versus AI as a transformative cognitive partner.

To fully categorize AI use and misuse, the researchers identified the specific academic integrity risks documented in the 35 studies and counted the mentions of academic integrity risks being discussed or focused on by the study. A total of 76 instances of integrity risks were identified in Table 2, indicating that many studies have uncovered multiple overlapping threats to academic integrity.

Table 2
Documented academic integrity risks

Integrity Risk	Nature	No. of Mentions
Loss of Critical Thinking/Over-reliance/ Dependency	Loss of intellectual effort, Loss of critical thinking, Decreased learning, Overdependence, or Excessive reliance	21
Plagiarism	AI-generated plagiarism, AI-assisted submissions, Ghostwriting, Copying, or directly substituting text	20
Cheating	Cheating, Dishonest Uses, AI Cheating/Deception, Academic misconduct, Misuse for cheating, or Academic dishonesty	18
Undisclosed Use / Non-compliance / Moral Hazard	Undisclosed AI Use, Non-declaration, Non-compliance, or Moral Hazard	4
Loss of Authorship/Authenticity/Creativity	Loss of authorial voice, Sacrificing creativity, Diminishing authenticity/originality, or Loss of artistic integrity	7
Other Technical/Specific Risks	Unauthorized resource use, Misinformation, Fabricated references, or Claiming false expertise	6

Categorization of AI Use in the SAMR Framework [19.1]. The researchers employed the SAMR framework (Substitution, Augmentation, Modification, and Redefinition) to categorize the integration of GenAI tools into student workflows. To ensure the reliability of the SAMR categorization, the researchers employed a collaborative, consensus-based approach (Levac et al., 2010) rather than independent coding. This allowed for the navigation of the "interpretive nuances" inherent in Generative AI use cases. The researchers also refine the data extraction matrix, establishing an iterative refinement process. A pilot extraction of five representative studies was conducted to establish baseline definitions for each SAMR level in the context of academic integrity. During the full extraction of the 35 studies, any "gray area" use cases (e.g., AI-driven paraphrasing versus AI-driven content generation) were flagged for synchronous discussion. These discussions led to the iterative refinement of the coding manual, ensuring that the boundaries between 'Enhancement' and 'Transformation' remained consistent across the dataset. A final audit was performed on the completed matrix to ensure that earlier extractions remained consistent with the refined definitions established in the later stages of the review.

Table 3
Documented SAMR level in GenAI context

SAMR Level	N	Percentage	SAMR Definition in GenAI Context
Substitution	8	22.9%	AI acts as a direct replacement for student effort with no functional change to the task.
Augmentation	16	45.7%	AI provides functional improvements (e.g., speed, grammar, or translation) to traditional tasks.
Modification	10	28.6%	The task is significantly redesigned around the AI's capabilities.
Redefinition	1	2.8%	AI enables entirely new academic tasks that were previously inconceivable.

As reflected in Table 3, analysis of the 35 studies revealed that the majority of current literature focuses on the Enhancement phase (68.6%), with a smaller portion exploring Transformative integration (31.4%). The distribution of the 35 studies across the four levels is summarized in Table 3 below. It is important to note that 31.4% (n=11) of the studies reported multi-level usage, where students utilized different AI functions within a single assignment.

To provide a holistic understanding of the current academic landscape, the researchers synthesized the three functional themes, their corresponding SAMR levels, and the documented integrity risks. This triangulation identifies specific "Risk Profiles" that exhibit a consistent thematic alignment where technological integration correlates with academic misconduct. Consistent with the framework of Arksey and O'Malley (2005) and the refinements of Levac et al. (2010), this study goes beyond descriptive counts to identify conceptual alignments between technology integration levels and integrity risks. This thematic synthesis does not seek to establish statistical correlation but rather to map the reported relationships within the current body of literature to inform the development of a prescriptive model.

Institutional Responses to Generative AI [1.1]. The next stage of the data synthesis involved the systematic gathering and compilation of institutional responses reported across the 35

studies. To ensure a data-driven result, the researchers utilized inductive thematic coding to group specific actions into six distinct themes based on their primary objective and implementation method.

A total of 45 response mentions were extracted (Table 4). The distribution shows a strong focus on formal governance and classroom-level modification, with a lower emphasis on purely technological solutions. Enforcement and updating institutional guidelines or policies emerged as the most prevalent response, indicating that institutions prioritize establishing regulations about Generative AI. These studies emphasized the importance of transparency, specifically through mandatory declaration requirements that require students to explicitly state whether and how AI was used. Interestingly, this category also included "total bans" in a minority of cases, though the trend moved toward "punishment clarity" and formalizing the definition of misconduct.

The second most frequent response focused on redesigning the assessment for students. Instead of prohibiting the tool, these responses altered the nature of the academic task. Common strategies included shifting to "viva voce" (oral exams), high-stakes in-class writing, or changing the grading rubric to reward the "human-in-the-loop" process (e.g., grading prompt histories or reflections) rather than the final generated product.

Despite the media focus on AI detectors as a technological intervention, this theme was notably less common in the empirical literature. Studies in this category explored the development and testing of detection algorithms or the use of secure testing environments (e.g., lockdown browsers). However, many studies that mentioned this theme also voiced skepticism on long-term efficacy of a purely technological "arms race."

A significant portion of the literature highlighted a gap in the AI literacy of institutions and faculty. Five studies reported a lack of institutional policy, leaving the burden of judgment on

individual instructors. This was often coupled with either faculty or curriculum development (n=4), which suggests that while some schools are not yet ready with rules, they are beginning to invest in the AI literacy of their staff and students through ethics modules and training.

Table 4
Institutional Response and Intervention

Institutional Action	Nature	No. of Mentions
Assessment Redesign/Modification	Changing assignment content, focus, format, or grading method	14
Policy/Guideline Update & Enforcement	Rules on use/banning, declaration requirements, supervision, punishment clarity, and formal policy updates	16
Technological Interventions	Development, testing, or implementation of GenAI detection tools or secure testing environments	5
Faculty Training/Curriculum Development	Introducing ethics modules, training for faculty, or mandatory student training/workshops	4
Ambiguity / Inaction / Clarification	Reports lack of policy, instructors relying on judgment, cautioning against detector reliance, or total bans	5
No Action Taken	Study was purely experimental with prescriptive outcomes	1

The SAMR-Integrity-Response Model. The final objective of this study was to propose an original prescriptive model. The researchers synthesized the functional themes, integration levels, and institutional actions from the 35 studies. This synthesis developed a framework that provides a strategic roadmap for educators to align their interventions with the specific nature of AI-assisted tasks.

The synthesis revealed three distinct "Risk-Response Profiles" based on the data:

1. **The Substitution Pattern.** When students engage in Generative Tasks, the AI functions at the Substitution level. Studies at this level are mostly distributed on Cheating/Misconduct (38.89%) and Plagiarism (25%). The primary response is Technological Intervention. This suggests that when AI is used as a direct replacement, institutions treat it as a technical security threat rather than a pedagogical one.

2. **The Augmentation/Modification Pattern.** For Linguistic Refinement, the AI acts as an Augmentation tool. Loss of Critical Thinking/Over-reliance is the most frequent risk (n=21), with 42.85% of these cases occurring in Augmentation and 38.09% in Modification. Despite this high cognitive risk, institutional responses here are often Ambiguous or limited to Policy Updates, leaving a "Pedagogical Gap."

3. **The Modification/Redefinition Pattern.** This is the most critical finding for modern education. When students use AI for Cognitive Support, it reaches the Modification or Redefinition level. At higher integration levels, the risk shifts to Loss of Authorship and Creativity (57.14% in Modification). This level identifies the highest need for Assessment Redesign, moving beyond simple rules to changing the nature of student output.

Based on these findings, the study proposes the SAMR-Integrity-Response (SIR) Framework (Table 5). This model provides a novel roadmap for mitigating AI-driven misconduct by aligning the intervention type with the specific risk profile of the integration level.

Table 5
The SAMR-Integrity-Response Model

SAMR Level	AI Use	Predominant Integrity Risk	Prescriptive Institutional Intervention
Substitution	Generative Task	Plagiarism & Direct Cheating	Robust detection, clear formative policies, and mandatory disclosure.
Augmentation	Linguistic Refinement	Over-reliance & Undisclosed Use	Explicit guidelines on "AI-assisted" vs. "AI-generated" work; citation standards.
Modification/Redefinition	Cognitive Support	Cognitive Erosion & Loss of Creativity	Shift to process-based grading, oral defenses, and in-class tasks. Mandatory AI ethics modules, faculty training, and co-creative frameworks.

The SIR Model addresses the mismatch found in the literature. While the data show that Policy Updates (n=16), Assessment Redesign (n=14), and Technological Intervention with AI

Detection (n=5) are common, they are often applied in a generic manner.

For "Enhancement" levels (Substitution & Augmentation), the model prescribes a defensive posture. Since the risk is the direct replacement of work (Plagiarism), technological and regulatory "guardrails" are necessary.

For "Transformative" levels (Modification & Redefinition), the model prescribes an evolutionary posture. Since the risk is the erosion of student thinking and authorial voice, detection is insufficient. Institutions must instead invest in Faculty Development and Assessment Redesign to evaluate the student's cognitive process rather than just final product.

DISCUSSION

The synthesis of 35 studies through the Integrated SIR (SAMR-Integrity-Response) Model reveals a fundamental shift in the academic landscape. The data confirm that as GenAI transitions from a tool for Generative Tasks and Linguistic Refinement to Cognitive Support, the nature of the "integrity threat" shifts from a breach of conduct to a potential erosion of student cognition.

This situation mirrors the historical resistance to previous technological disruptions, such as the introduction of the handheld calculator or the transition to internet-based search engines. However, it is essential to acknowledge that GenAI has created a more profound shift by generating entirely new outputs through predictive modeling.

While the majority of the included studies focus on Higher Education, the SIR model's logic is generalizable across the contexts of K-12 education and professional training, with minor adjustments. In the context of K-12 education, the risk of Cognitive Erosion is higher, as students are still developing foundational "lower-order" skills. Interventions should lean toward Policy and Technical guardrails to protect the developmental phase. The timing of

AI integration in education is a critical factor in the cognitive development of individuals. Overreliance on AI-generated suggestions during adolescence can limit neural development in the prefrontal cortex, which is essential for independent reasoning and problem-solving (Riser, 2025). Because this developmental phase is vital for identity formation and intellectual growth, it requires deliberate practice rather than the shortcuts that are offered by automation.

Consequently, a developmental threshold must be established before introducing such tools to students. This is supported by Park and Milner (2025), whose research on faculty perspectives suggests that advanced AI tools are best suited for the university level; at this stage, students are presumed to have already solidified the foundational cognitive abilities and critical thinking habits that were nurtured throughout their primary and secondary education.

A primary finding of this study is the mismatch in current institutional actions. While the results indicate that Cognitive Erosion/Over-reliance is the most frequent risk (n = 21), institutional responses remain heavily weighted toward Policy/Guideline Enforcement (n = 16) and Technological Interventions (n = 5), specifically AI Detection Tools. This suggests that many institutions are attempting to solve "Process-oriented" risks (Modification level) with "Product-oriented" solutions (Substitution level). As indicated by the SIR Model, a policy-only approach is insufficient for the Cognitive Support theme; instead, a pedagogical shift toward assessment redesign is required to protect the "human-in-the-loop" requirement of learning.

To operationalize the SIR (SAMR-Integrity-Response) Model, institutions must move beyond reactive "policing" and toward proactive "pedagogical transformation." Institutions should follow a three-phase timeline:

Phase 1: Policy Synchronization. Immediate update of honor codes to define "AI-Assisted" vs. "AI-Generated."

Phase 2: Pedagogical Pivot. Training faculty in Assessment Redesign to move toward oral exams and process-based grading.

Phase 3: Curricular Integration. Making "AI Literacy" a mandatory ethics module for all students (Redefinition level).

To implement a comprehensive response to GenAI, institutions should begin with Phase 1, which is the *Policy Synchronization*. Recent literature highlights that ambiguity in institutional guidelines is the leading predictor of undisclosed AI use and unintentional misconduct (Hutson, 2024). Because students often perceive AI as a simple refinement tool, institutions must immediately establish a "ground truth" through honor codes. This phase involves the immediate update of honor codes to eliminate the "moral hazard" created by ambiguous regulations. Honor codes must transition from "punitive" to "descriptive." By explicitly defining "AI-Assisted" (human-led refinement) vs. "AI-Generated" (AI-led creation), institutions provide students with the "epistemic vigilance" needed to navigate the Substitution tier (Lund et al., 2025)

The transition continues with Phase 2, which is the *Pedagogical Pivot*. This is where educators are trained in Assessment Redesign to move toward "viva voce" (oral exams) and process-based grading. This shift moves the focus from a final static product to evaluating the "learning journey." By examining the various stages of a student's work, such as initial drafts, revisions, and reflective journals, educators can prioritize the quality of iterative improvements over the end result. (Wang, 2024) Adopting Project-Based Learning (PBL) is also central to this phase, as research shows that students who use PBL achieve higher academic performance and significantly lower similarity percentages compared to those completing traditional assignments. These authentic assessments, which include live demos and practical examinations, are the most resilient against AI-facilitated cheating because they require students to demonstrate their reasoning in real-time. (Kldiashvili et al., 2025).

While these pedagogical shifts ensure a more authentic evaluation, their success will depend on equitable access to tools. However, the synthesis also highlights a growing concern regarding integrity and equity. Access to advanced LLMs is increasingly behind a "paywall". Students with access to premium models may produce higher-quality work than those using free, less-capable models. As Vesna et al. (2025) observe, the digital divide in AI-driven education is fueled by technological infrastructure gaps and socio-economic barriers that disproportionately impact marginalized communities and underfunded institutions, which creates an environment where the ability to benefit from AI-enhanced learning is restricted to those within well-funded ecosystems. This shifts the burden of 'integrity' onto the individual's financial capacity rather than their moral character." This creates a moral hazard where integrity is no longer just about honesty, but about socio-economic access. Supporting this, Smit et al. (2025) emphasize that student confidence is a critical factor in university policy, as allowing or banning these technologies can disadvantage specific groups, particularly those unable to afford or access paid versions of advanced AI tools. Consequently, institutions must consider how Policy Updates can account for this disparity, perhaps by providing campus-wide access to specific tools to level the playing field.

Lastly is Phase 3, *Curricular Integration*. This phase establishes "AI Literacy" as a mandatory ethics module for all students, operating at the Redefinition level. It must go beyond introducing basic prompting to include digital ethics, data privacy, and the critical evaluation of algorithmic bias. (Riser, 2025) Students must also be trained in structured frameworks, such as Lo's CLEAR model (Concise, Logical, Explicit, Adaptive, Reflective) and the DEER praxis (Define, Evaluate, Explore, Reflect), to ensure they engage with AI as reflective partners rather than authoritative sources. (Cummings et al., 2024) Ultimately, this integration transforms the teacher into an "AI ethics coach" who facilitates reflexivity, empowering students to maintain their unique authorial voice and

critical independence in an increasingly automated world.

Conclusion. This research sought to systematically identify and categorize the ways in which the literature reports the uses of generative AI among high school and college students, and determine the frequency distribution across SAMR levels along with the specific types of institutional responses concerning the use of generative AI to propose an original prescriptive model that maps integrity risks to institutional interventions across the SAMR levels, establishing a novel framework for mitigating AI-driven academic misconduct.

The findings indicate that most studies on Generative AI and academic integrity are primarily focused on higher education institutions, such as colleges and universities, with ChatGPT being the leading generative artificial intelligence tool for students. Three themes were identified in classifying the way students use generative AI: for generative tasks, linguistic refinement, and as cognitive support. These themes were identified according to how the AI tools function within the academic workflow of the students. Students often misuse generative AI through over-reliance, thereby losing their critical thinking skills, and through the use of AI-generated or AI-assisted plagiarism. Furthermore, in categorizing the integration of GenAI tools using the SAMR model, the majority of the studies and misuse are categorized on the Enhancement phase, specifically on the Augmentation level. Aside from the SAMR categorization, the institutional responses and interventions mentioned in the studies are also collected. It was determined that updating and enforcing institutional policies and guidelines, as well as assessment and pedagogical redesigns, are the common interventions implemented. Synthesizing the themes formulated on the student use of GenAI, the categories of technology integration in the SAMR model, and the institutional responses and interventions, an original prescriptive model was formulated that provides a strategic

roadmap for educators to align their interventions with the specific nature of AI-assisted tasks. This original framework offers a new lens for examining the various unethical uses of Generative AI in education, following the three phases of implementing the SAMR-Integrity-Response Model. By formalizing the original framework, which is the SAMR-Integrity-Response Model, this study provides a comprehensive response to the various unethical uses of GenAI in education.

This scoping study is subject to several limitations. First, the search was limited to specific academic databases, potentially missing gray literature or white papers from industry leaders. Second is the ongoing advancement in GenAI technology. GenAI is evolving faster than the peer-review cycle; some studies published in 2023 may not reflect the capabilities of 2025 models. Lastly, many studies were pilot programs with small sample sizes, which may limit the broad applicability of some reported Institutional Actions. To address these gaps, future research should adopt more agile and cross-sectoral methodologies, such as living systematic reviews, where findings are updated in real-time. This is essential to keep pace with AI systems that shift from passive assistants to autonomous agents. Future studies must also test frameworks like the SIR Model across multi-institutional and global contexts. This will determine if specific interventions can be effectively scaled to diverse socio-economic groups. Finally, a critical new direction for research is the investigation of corporate responsibility and algorithmic accountability, with a focus on the role of AI developers in ensuring educational equity.

Author contributions. (Not available)

Conflict of interest. The authors declare no conflict of interest.

Funding source. This research received no external funding.

Artificial intelligence use. No AI tools were used in the preparation of this manuscript.

Ethics approval statement. Ethics approval was not required for this study as it involved publicly available data.

Data availability statement. All data supporting the findings of this study are included within the manuscript and its supplementary materials.

Acknowledgement. (Not available)

Publisher's disclaimer. The views expressed in this article are those of the authors and do not necessarily reflect the views of the publisher. The publisher disclaims any responsibility for errors or omissions.

REFERENCES

- Akiba, D., & Garte, R. (2024). Leveraging AI tools in university writing instruction: Enhancing student success while upholding academic integrity. *Journal of Interactive Learning Research*, 35(4), 467-480.
- Alawad, E. A., Ayadi, H. H., & Alhinai, A. A. (2025). Guarding integrity: A case study on tackling AI-Generated Content and Plagiarism in academic writing. *Theory and Practice in Language Studies*, 15(6), 1730-1742.
- Arias Ortiz, E., Castro, N., Forero, T., Gambi, G., Giambruno, C., Pérez-Alfaro, M., Rodríguez Segura, D. (2025). *AI and education: Building the future through digital transformation (Technical Note No. IDB-TN-03122)*. Inter-American Development Bank.
- Arksey, H., & O'Malley, L. (2005). Scoping studies: Towards a Methodological Framework. *International Journal of Social Research Methodology*, 8(1), 19-32.
- Aromataris, E., Lockwood, C., Porritt, K., Pilla, B., & Jordan, Z. (Eds.). (2024). *JBI manual for evidence synthesis*. JBI.
- Balalle, H., & Pannilage, S. (2025). Reassessing academic integrity in the age of AI: A systematic literature review on AI and academic integrity. *Social Sciences & Humanities Open*, 11(1), 101299.
- Blundell, C. N., Mukherjee, M., & Nykvist, S. (2022). A scoping review of the application of the SAMR model in research. *Computers and Education Open*, 3, 100093. <https://doi.org/10.1016/j.caeo.2022.100093>
- Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3(2), pp. 77-101. ISSN 1478-0887
- Bretag, T., Harper, R., Burton, M., Ellis, C., Newton, P., & van Haeringen, K. (2019). Contract cheating and assessment design: Exploring the relationship. *Assessment & Evaluation in Higher Education*, 44(5), 676-691.
- Caling, J. G., Antonio, J. K. T., Dimatatac, M. F. L., Sabellano, M. D., Vicencio, V. D. R., Prias, J. M., & Ramos, J. C. M. (2025). Pre-service teachers' use of ChatGPT and acquired moral dissonance. *Journal of Interdisciplinary Studies in Education*, 14(4), 73-98.
- Cummings, R. E., Monroe, S. M., & Watkins, M. (2024). Generative AI in first-year writing: An early analysis of affordances, limitations, and a framework for the future. *Computers and Composition*, 71, Article 102827.
- Draxler, F., Werner, A., Lehmann, F., Hoppe, M., Schmidt, A., Buschek, D., & Welsch, R. (2023). The AI Ghostwriter Effect: When Users Do Not Perceive Ownership of AI-Generated Text But Self-Declare as Authors. *ACM Transactions on*

- Computer-Human Interaction*, 31(2), 1-40. Article 25.
- Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., Baabdullah, A. M., Koochang, A., Raghavan, V., & Ahuja, M. (2023). "So what if ChatGPT wrote it?" Multidisciplinary perspectives on opportunities, challenges, and implications of generative conversational AI for research, practice, and policy. *International Journal of Information Management*, 71, 102642.
- Ellis, C., Barlow, L., Morris, J., & Van Haeringen, K. (2023). Where does the human stop and the machine begin? Academic integrity, human-machine collaboration, and the rise of Generative AI. *Journal of Academic Ethics*, 21(3), 299–318.
- Hutson, J. (2024). Rethinking Plagiarism in the Era of Generative AI. *Journal of Intelligent Communication*, 3(2), 20–31.
- Gruenhagen, J. H., Sinclair, P. M., Carroll, J. A., Baker, P. R. A., Wilson, A., & Demant, D. (2024). The rapid rise of generative AI and its implications for academic integrity: Students' perceptions and use of chatbots for assistance with assessments. *Computers and Education: Artificial Intelligence*, 7, 100273.
- International Center for Academic Integrity [ICAI]. (2021). *The Fundamental Values of Academic Integrity*. (3rd ed.).
- Kldiashvili, E. D., Mamiseishvili, A., & Zarnadze, M. (2025). Academic integrity within the medical curriculum in the age of generative artificial intelligence. *Health Sci Rep.*, 8(2):e70489. <https://doi.org/10.1002/hsr2.70489>.
- Kofinas, A. K., Tsay, C.-H., & Pike, D. (2025). The impact of generative AI on academic integrity of authentic assessments within a higher education context. *British Journal of Educational Technology*, 56, 2522-2549. <https://doi.org/10.1111/bjet.13585>
- Landis, J. R., & Koch, G. G. (1977). The measurement of observer agreement for categorical data. *Biometrics*, 33(1), 159–174.
- Levac, D., Colquhoun, H., & O'Brien, K. K. (2010). Scoping studies: Advancing the methodology. *Implementation Science*, 5(1), 1-9. <https://doi.org/10.1186/1748-5908-5-69>
- Lund, B. D., Wang, T., & Hutson, J. (2025). From punitive to descriptive: Transitioning honor codes for the generative AI era. *Education Sciences*, 15(1), 2.
- Morris, M. R., Sohl-Dickstein, J., Fiedel, N., Warkentin, T., Dafoe, A., Jernite, Y., Farquhar, V., & Legg, S. (2023). Levels of AGI: Operationalizing progress on the path to AGI. *arXiv*.
- Park, J. J., & Milner, P. (2025). Enhancing academic writing integrity: Ethical implementation of generative artificial intelligence for non-traditional online students. *TechTrends*, 69(1), 176–188.
- Puentedura, R. R. (2006). *Transformation, technology, and education*. Hippasus. <http://www.hippasus.com/resources/tte/>
- Smit, M., Bond-Barnard, T., & Wagner, R. F. (2025). Artificial intelligence in South African higher education: Survey data of master's level students. *Project Leadership & Society*, 6, 100142.
- Ratten, V. (2023). The post COVID-19 pandemic era: Changes in teaching and learning methods for management educators. *The International Journal of Management Education*, 21(2), Article 100777.

- Rehman, Z. U., & Aurangzeb, W. (2021). The SAMR Model and Bloom's Taxonomy as a Framework for Evaluating Technology Integration at University Level. *Global Educational Studies Review*, VI(IV), 1-11.
- Riser. (2025). Generative Artificial Intelligence (AI) in the high school English classroom and its effect on students' identity and teachers' practice. *English in Education*, 59(2), 176-192.
- Roe, J., Perkins, M., Chonu, G. K., & Bhati, A. (2023). Student perceptions of peer cheating behaviour during COVID-19 induced online teaching and assessment. *Journal of Higher Education Policy and Management*, 44(6), 612-627. <https://doi.org/10.1080/07294360.2023.2258820>
- Tricco, AC, Lillie, E, Zarin, W, O'Brien, KK, Colquhoun, H, Levac, D, Moher, D, Peters, MD, Horsley, T, Weeks, L, Hempel, S et al. PRISMA extension for scoping reviews (PRISMA-ScR): checklist and explanation. *Ann Intern Med.*, 169(7):467-473.
- UNESCO. (2023). *Global education monitoring report 2023: Technology in education: A tool on whose terms?* <https://www.unesco.org/en/articles/global-education-monitoring-report-2023-technology-education-tool-whose-terms>
- University Center for Teaching and Learning [UCTL]. (n.d.). *What is generative AI?* University of Pittsburgh. <https://teaching.pitt.edu/resources/what-is-generative-ai/>
- Vesna, L., Sawale, P. S., Kaul, P., Pal, S., & Murthy, B. S. N. V. R. (2025). Digital divide in AI-powered education: Challenges and solutions for equitable learning. *Journal of Information Systems Engineering & Management*, 10(21s), 300-308.
- Waltzer, T., Pilegard, C., & Heyman, G. D. (2024). Can you spot the bot? Identifying AI-generated writing in college essays. *International Journal for Educational Integrity*, 20(1), 1-18. <https://doi.org/10.1007/s40979-024-00158-3>
- Wang, C. (2024). Exploring students' generative AI-assisted writing processes: Perceptions and experiences from native and nonnative English speakers. *Technology, Knowledge and Learning*, 30(1), 1825-1846.
- Wang, S., Wang, F., Zhu, Z., Wang, J., Tran, T., & Du, Z. (2024). Artificial intelligence in education: A systematic literature review. *Expert Systems with Applications*, 252, Article 124167. <https://doi.org/10.1016/j.eswa.2024.124167>