

A theoretical model for the ethical use of generative AI tools in higher education assessments

International
Journal of Ethics
and Systems

Maria Ijaz Baig 

*Applied Science Research Center, Applied Science Private University,
Amman, Jordan, and*

Elaheh Yadegaridehkordi 

*School of Engineering and Technology, College of Information and
Communications Technology, CQUniversity, Sydney, Australia*

Received 22 September 2025

Revised 12 January 2026

13 March 2026

13 May 2026

4 July 2026

Accepted 5 July 2026

Abstract

Purpose – This study aims to develop a theoretical framework to identify the key factors influencing the ethical use of generative artificial intelligence (GenAI) tools in academic assessments from a student perspective.

Design/methodology/approach – Grounded in digital literacy theory and self-determination theory (SDT), the study examines six constructs: GenAI transparency, GenAI autonomy, GenAI privacy, GenAI ethics literacy, GenAI policy clarity and GenAI fairness. Data was collected from 171 university students through an online survey and analyzed using partial least squares structural equation modeling (PLS-SEM).

Findings – The results show that GenAI transparency, GenAI privacy, GenAI ethics literacy, GenAI policy clarity and GenAI fairness significantly influence the ethical use of GenAI tools in assessments, while GenAI autonomy did not have a significant effect.

Practical implications – The findings provide actionable insights for students, educators, institutions and developers to promote the ethical use of GenAI tools in academic settings. Establishing clear policies, strengthening ethics literacy and ensuring fairness and transparency can foster responsible integration of GenAI in assessments.

Originality/value – This study is original in its focus on the ethical use of GenAI in academic assessments from a student perspective, a dimension largely overlooked in prior research. By integrating under-explored constructs such as policy clarity, fairness and ethics literacy, the study contributes a novel theoretical framework and empirical evidence to guide students' responsible use of GenAI and inform institutional practices for ethical GenAI adoption.

Keywords Generative artificial intelligence (GenAI), Ethical use, Higher education, Assessments, Digital literacy, Self-determination theory

Paper type Research paper

1. Introduction

The adoption of artificial intelligence (AI) across industries has significantly transformed multiple domains, including education (Law, 2024; Zhuang *et al.*, 2025). AI supports personalized learning, reduces routine administrative work and improves instructional effectiveness (Yim and Su, 2025). Among AI advancements, generative artificial intelligence



© Maria Ijaz Baig and Elaheh Yadegaridehkordi. Published by Emerald Publishing Limited. This article is published under the Creative Commons Attribution (CC BY 4.0) licence. Anyone may reproduce, distribute, translate and create derivative works of this article (for both commercial and non-commercial purposes), subject to full attribution to the original publication and authors. The full terms of this licence may be seen at <http://creativecommons.org/licences/by/4.0/>

International Journal of Ethics and
Systems
Emerald Publishing Limited
2514-9369
DOI 10.1108/IJOES-09-2025-0538

(GenAI) tools have emerged as particularly transformative, enabling content creation, complex task automation and decision-making support through the processing of vast data sets (Nguyen, 2025; Agustini, 2023). ChatGPT remains the most widely used GenAI tool in education, while Bing AI, Copilot, OpenAI Playground and Bard AI are gaining traction (Foung *et al.*, 2024). GenAI has been applied in personalized learning, content creation, feedback delivery and assessment, demonstrating potential to enhance educational outcomes (Oved and Alt, 2025).

However, the integration of GenAI into academic assessments introduces ethical challenges. Responsible use requires students to engage with these tools transparently, comply with academic integrity policies, avoid plagiarism or misrepresentation and critically evaluate outputs (Beynen, 2024; Zhang and Magerko, 2025). Fairness must be ensured so that AI does not provide undue advantage or disadvantage and tools should support, not replace, authentic learning (Zlotnikova *et al.*, 2025). Institutional support, reflective thinking and adherence to ethical standards are essential for responsible AI engagement in assessment contexts (Holmes *et al.*, 2022; Đerić *et al.*, 2025). As GenAI becomes increasingly embedded in education, its use in assessments raises concerns regarding ethical conduct and overreliance (Izevbigie *et al.*, 2025; Turner, 2025). Higher education institutions across the globe are developing regulatory policies to govern students' use of such tools during assessments (Luo, 2024). While these tools can enhance effectiveness and equity in academic processes, they also introduce challenges related to bias, privacy and transparency (Memarian and Doleck, 2023). Assessment outcomes directly affect students' academic progress (Luo, 2025). Therefore, safeguarding integrity and fairness requires a commitment to the responsible application of these technologies.

While studies have examined ethical AI use in higher education generally (Baig and Yadegaridehkordi, 2025; Đerić *et al.*, 2025; Izevbigie *et al.*, 2025), research focusing specifically on GenAI in assessments remains limited. Existing studies often foreground educators' perspectives (Khlaif *et al.*, 2025; Bower *et al.*, 2024) or compare student and educator perceptions (Roe *et al.*, 2024). Although foundational AI ethics research outlines principles such as transparency, privacy, accountability and fairness (Huang, 2023; Shin, 2021), there is no theory-informed framework that operationalizes these principles for students' ethical use of GenAI in assessments. Without such guidance, students, as the primary end users, lack structured support to make responsible choices regarding disclosure, attribution and appropriate assistance. A structured framework would also inform policy formulation that fosters responsible innovation in assessment practices (Shuford, 2024; Adarkwah, 2025; García-López and Trujillo-Liñán, 2025). To address this gap, this study combines digital literacy theory (Gilster, 1997) and self-determination theory (SDT) (Deci and Ryan, 1985) to propose and empirically validate a model for the ethical use of GenAI tools by students in academic assessments. This model translates ethical principles into measurable constructs and links students' competencies and motivational factors to their ethical decision-making in assessment contexts.

This paper is structured as follows: "Section 2 presents the literature review; Section 3 outlines the theoretical foundation and hypothesis development; Section 4 describes the research methodology, including data collection and respondent details; Section 5 reports the results; Section 6 provides the discussion; and Section 7 outlines the implications for theory and practice."

2. Literature review

The integration of GenAI into higher education assessment has attracted growing scholarly attention due to its transformative potential and associated ethical challenges. While existing

studies provide valuable insights, the literature remains fragmented and largely descriptive. A closer examination reveals that prior research is shaped by several underlying theoretical tensions. Accordingly, this section synthesizes the literature by organizing it around key debates to provide a more analytical and theoretically grounded foundation for understanding ethical GenAI use in assessment contexts. [Table 1](#) summarizes GenAI educational-related studies in the assessment context.

One of the most prominent debates in literature concerns the tension between leveraging GenAI for innovation and maintaining academic integrity. On the one hand, GenAI is widely recognized as a transformative tool capable of enhancing learning, improving feedback quality and supporting the redesign of assessments to promote higher-order thinking skills ([Furze et al., 2024](#); [Khlaif et al., 2025](#)). For example, [Furze et al. \(2024\)](#) introduced the artificial intelligence assessment scale, which provides a structured approach to integrating AI into assessment at varying levels. Similarly, [Khlaif et al. \(2025\)](#) proposed an adaptive framework that encourages educators to explore and adopt GenAI to support more flexible and responsive assessment practices.

On the other hand, this rapid integration raises critical concerns regarding academic integrity, authorship and the validity of assessment outcomes. Prior studies highlight risks such as overreliance on AI-generated outputs, challenges in distinguishing between human and AI contributions and concerns about fairness and transparency in evaluation processes ([Eaton et al., 2025](#); [Corbin et al., 2025](#); [Jukiewicz, 2024](#)).

While GenAI can enhance efficiency and consistency, for example, in automated grading, it also introduces issues such as hallucinations and a lack of reproducibility, limiting its reliability in high-stakes contexts ([Jukiewicz, 2024](#)).

A second debate concerns the distinction between technical and ethical competence in GenAI use. Much of the existing literature focuses on individuals' ability to effectively use AI tools, including skills related to prompting, tool selection and adaptation ([Foung et al., 2024](#); [Bower et al., 2024](#)). These studies highlight the importance of digital literacy in enabling users to engage with GenAI technologies productively. However, relatively less attention has been given to ethical competence, including the knowledge, values and decision-making processes required for responsible AI use. For instance, [Bower et al. \(2024\)](#) examined educators' perceptions of how teaching and assessment should evolve in response to GenAI, identifying the importance of ethical awareness alongside technical skills. [Roe et al. \(2024\)](#) demonstrated that both students and staff express concerns about academic integrity and the reliability of AI tools, emphasizing the need for clear policies and guidance. Nevertheless, these insights remain fragmented and are not integrated into a broader theoretical framework.

This distinction suggests that technical proficiency alone is insufficient to ensure ethical behavior. Existing studies do not adequately explain how users translate technical capabilities into ethically responsible actions or how ethical considerations are internalized and applied in practice. Consequently, there is a lack of theoretical clarity regarding the relationship between competence and ethical behavior in GenAI use.

Taken together, these debates highlight that ethical GenAI use cannot be explained solely by technological capability or institutional regulation. Instead, responsible behavior emerges from the interplay between competence, motivation and governance mechanisms. However, the current literature remains fragmented, as these dimensions are typically examined in isolation rather than within an integrated theoretical framework. Moreover, much of the prior research is either descriptive or practice-oriented, focusing on perceptions, challenges and implementation strategies without offering a structured explanation of behavioral processes. As a result, there is limited understanding of how students translate ethical knowledge into

Table 1. GenAI educational-related studies in the assessment context

Author	Context	Data collection	Data analysis	Constructs/key findings
Khlaif <i>et al.</i> (2025)	Framework for redesigning assessments with gen AI-enhanced learning	Semi-structured interviews and four focus group sessions with 61 participants	Inductive thematic analysis	"Need for adaptive assessment framework, focus on fairness and educator readiness"
Eaton <i>et al.</i> (2025)	Focusing on the integration of GenAI chatbots in assessments	Interviews with 28 higher education staff members	Thematic narrative analysis of the interviews	"Navigated ethical boundaries for chatbot use, supporting integration for prompting and critical thinking skills"
Corbin <i>et al.</i> (2025)	Focus on the challenges of assessment design in the era of GenAI	Semi-structured individual interviews with 20 teachers	Applied Rittel and Webber's framework of wicked problems to examine the interview data	"Found GenAI-related assessment challenges exhibit all ten characteristics of wicked problems"
Furze <i>et al.</i> (2024)	Pilot implementation of GenAI-supported assessment scale (AIAS)	Institutional records and academic staff report	Descriptive analysis	"Highlighted the potential integration of GenAI into higher education assessment"
Bower <i>et al.</i> (2024)	Educators' perceptions of the impact of GenAI on teaching and assessment	Survey 318 educators	Thematic analysis of responses	"Perceived impact of GenAI, motivation to change teaching and assessment, curriculum and pedagogy shifts (e.g. learning with AI, higher-order thinking, ethics); influence of demographics and AI awareness"
Foung <i>et al.</i> (2024)	Redesign of assessments in an elective English course	'Written reflections from 74 students; Focus group interviews with 28 students'	Analysis of student reflections and interviews	"Tool selection skills, flexibility in AI tool use; equitable access to AI tools, student perceptions and concerns regarding AI tool features"
Jukiewicz (2024)	Evaluation of ChatGPT as an auto-grading tool	Grades from nine programming assignments assessed by both the teacher and ChatGPT	Comparative grade analysis	"Measured grading consistency"
Roe <i>et al.</i> (2024)	Perspectives of students and educators on the use of GenAI in assessment and feedback	Survey 35 academic staff and 282 students	Descriptive statistics and reflexive thematic analysis	"Concerns about academic integrity, usefulness of AI feedback"

responsible GenAI use within assessment contexts. Accordingly, this study addresses this gap by developing an integrated theoretical framework grounded in digital literacy theory and SDT, while incorporating key ethical governance dimensions identified in prior research. This approach responds directly to the identified debates by providing a unified explanation of how competence, motivation and governance interact to shape ethical GenAI use in higher education assessment.

3. Theoretical foundations

Digital literacy provides the competency foundation for understanding and navigating digital environments. Originally conceptualized as the ability to access and evaluate digital information (Gilster, 1997), the framework has evolved to encompass critical thinking, ethical awareness, privacy management and responsible decision-making (Eshet, 2004; Ng, 2012). In AI-mediated academic contexts, digital literacy includes the capacity to evaluate algorithmic outputs, recognize potential biases, understand data protection implications and interpret institutional regulations governing technology use. Ethics literacy, identified as a central dimension (Belshaw, 2012), reflects learners' ability to apply ethical reasoning when engaging with digital tools. In GenAI-assisted assessments, digital literacy therefore represents the cognitive and evaluative competence necessary for informed and responsible use.

While digital literacy explains students' capacity to evaluate information and identify ethical issues, it does not explain why they choose to act ethically. SDT addresses this motivational dimension by proposing that autonomy promotes the internalization of ethical standards and supports self-regulated behavior (Deci and Ryan, 1985; Ryan and Deci, 2000). In academic settings, autonomy facilitates the internalization of externally imposed rules and standards. When students perceive themselves as autonomous decision-makers, ethical principles are more likely to be self-endorsed and consistently enacted. Autonomy is therefore conceptualized as the central motivational mechanism underlying ethical engagement with GenAI tools. Prior research suggests that competence alone is insufficient to drive behavior unless it is supported by intrinsic or autonomous motivation (Ryan and Moller, 2017). Therefore, the ethical use of GenAI can be understood as emerging from the interaction between competence (digital literacy) and motivation (autonomy). Students are more likely to engage in responsible and transparent AI use when they not only possess the necessary literacy but also experience a sense of autonomy in their decision-making. In this way, autonomy facilitates the translation of digital competence into ethical behavior, providing a clear behavioral mechanism that links knowledge to action.

By integrating digital literacy and autonomy, this study explains how ethical GenAI use emerges through the interaction of competence and motivation. Digital literacy equips students with the cognitive capacity to recognize ethical issues such as bias, data privacy risks and inappropriate reliance on AI outputs and to critically evaluate alternative courses of action. However, the ability to recognize and evaluate ethical dilemmas does not in itself ensure ethical behavior. Drawing on SDT, autonomy explains how this competence is translated into action. When students experience a sense of volition and psychological ownership, they are more likely to internalize ethical standards rather than comply with them superficially. This internalization process transforms externally imposed rules into self-endorsed principles, increasing the likelihood of consistent and principled decision-making. Accordingly, digital literacy shapes ethical awareness and judgment, while autonomy drives the internalization and enactment of those judgments. Ethical GenAI use therefore emerges through a process in which competence supports ethical judgment and autonomy facilitates the enactment of ethical choices.

Thus, this study operationalizes core ethical principles into six constructs that reflect conditions relevant to the use of GenAI in academic assessments. While the proposed model is primarily grounded in digital literacy theory and SDT, these constructs are also supported by Huang's (2023) normative ethics framework, which provides complementary justification for the selected ethical dimensions. Specifically, GenAI transparency refers to the explainability of AI outputs and users' ability to justify their application. GenAI fairness captures equitable treatment and the avoidance of discrimination or undue academic advantage. GenAI privacy concerns the protection of personal data and informational rights in AI-mediated environments. GenAI policy clarity reflects the extent to which institutional guidelines clearly define acceptable AI practices and associated accountability structures. GenAI ethics literacy represents students' awareness and understanding of ethical considerations in AI-assisted academic work. Finally, GenAI autonomy reflects students' perceived agency and control in decision-making when engaging with GenAI tools.

3.1 Hypotheses

GenAI transparency refers to the extent to which users receive clear, accessible and accurate information about the system's functioning and content generation mechanisms (Luo, 2025). It is conceptually distinct from privacy, which concerns data protection, while transparency concerns rule visibility and interpretability (Dawson, 2020; Huang *et al.*, 2024). A higher level of transparency allows learners to gain clearer insights into the strengths and constraints of GenAI tools, thereby fostering responsible and ethical use (Rasul *et al.*, 2024). Transparent systems can help promote trust, increase user accountability and reduce the likelihood of misuse, such as plagiarism or misrepresentation of AI-generated work (Thulasiram, 2025):

H1. GenAI transparency is positively related to ethical GenAI use in assessments.

GenAI autonomy pertains to the degree of control and discretion users possess over the utilization of GenAI tools within academic settings (Giannakos *et al.*, 2024). It reflects responsible decision-making under institutional constraints and does not imply unrestricted use (He, 2025). A high level of perceived autonomy can enable learners to independently regulate their use of these technologies, integrating them as supportive resources while maintaining ownership of their academic work (Agustini, 2023). This sense of ownership encourages thoughtful choices, as students are more inclined to critically evaluate AI-generated content to ensure it aligns with academic standards and ethical principles (Nguyen, 2025). Institutions that promote autonomous engagement with GenAI tools can cultivate an environment where ethical considerations are integral to the learning process (Szabó and Szoke, 2024):

H2. GenAI autonomy is positively related to ethical GenAI use in assessments.

GenAI privacy refers to the protection of personal data and the confidentiality of students' interactions with AI systems (Shailendra *et al.*, 2024). Confidence in data protection makes it more likely that students will use GenAI tools responsibly, in ways that respect privacy regulations and uphold academic integrity (Baig and Yadegaridehkordi, 2025). In the context of academic assessments, strong privacy protections help reduce the risk of unethical behaviors, such as unauthorized sharing of sensitive information, cheating or plagiarism (Dawson, 2020). Safeguarding user privacy thus plays a critical role in promoting responsible and ethical use of GenAI tools (Huang *et al.*, 2024). Therefore, ensuring privacy

in GenAI tools can encourage ethical use, as students are more likely to follow academic guidelines when they feel that their data are secure (Law, 2024):

H3. GenAI privacy is positively related to ethical GenAI use in assessments.

GenAI ethics literacy refers to the understanding of the ethical implications of using AI technologies (Al-kfairy *et al.*, 2024). Students educated about the ethical challenges and responsibilities associated with GenAI are more likely to use the tools in a manner that adheres to academic integrity standards (Corbin *et al.*, 2025). Educating students about responsible AI use, including avoiding plagiarism and ensuring fairness, can promote ethical decision-making in academic assessments (Chan, 2023). Ethics literacy can be important in fostering ethical GenAI use in assessments, as it equips students with the knowledge to make responsible decisions (Xia *et al.*, 2024):

H4. GenAI ethics literacy is positively related to ethical GenAI use in assessments.

GenAI policy clarity involves clear communication from educational institutions regarding the rules and guidelines for using GenAI tools in assessments (Luo, 2024). Universities that establish well-defined policies increase the likelihood that students will understand what constitutes ethical use of the technology (Moorhouse *et al.*, 2023). Clear guidelines help prevent academic misconduct and ensure that students use GenAI tools in a manner consistent with academic integrity standards (Malik *et al.*, 2023). Therefore, providing clear policies on GenAI use can positively influence its ethical use in assessments by guiding students to adhere to institutional expectations (Law, 2024):

H5. GenAI policy clarity is positively related to ethical GenAI use in assessments.

GenAI fairness refers to the equitable treatment of all users, ensuring that AI systems are free from bias and discrimination (Shuford, 2024). In academic assessments, fairness ensures that all students have equal access to GenAI tools and are not unfairly advantaged or disadvantaged by the technology (Dabis, and Csáki, 2024). Students who perceive GenAI tools as fair are more likely to use them ethically, understanding that the tools provide an equal opportunity for success (Holmes *et al.*, 2022). Therefore, fairness in GenAI can positively influence ethical use in assessments by promoting equitable and unbiased outcomes (Rasul *et al.*, 2024):

H6. GenAI fairness is positively related to ethical GenAI use in assessments.

This model can also be viewed within a hierarchical construct architecture where GenAI policy clarity and GenAI fairness represent governance drivers, GenAI transparency, GenAI privacy and GenAI ethical awareness represents cognitive drivers, GenAI autonomy reflects motivational drivers and ethical GenAI use represents behavioral outcomes. The proposed research model is presented in Figure 1.

4. Methodology

This study seeks to propose a theoretical framework that addresses the ethical application of GenAI tools by students in academic assessments. To evaluate the research hypotheses, a quantitative design was used and partial least squares structural equation modeling (PLS-SEM) was selected due to its strong predictive capabilities and suitability for analyzing complex models (Hair *et al.*, 2019). The data were examined using SmartPLS version 3.0. PLS-SEM was also chosen because it is particularly appropriate for theory development of

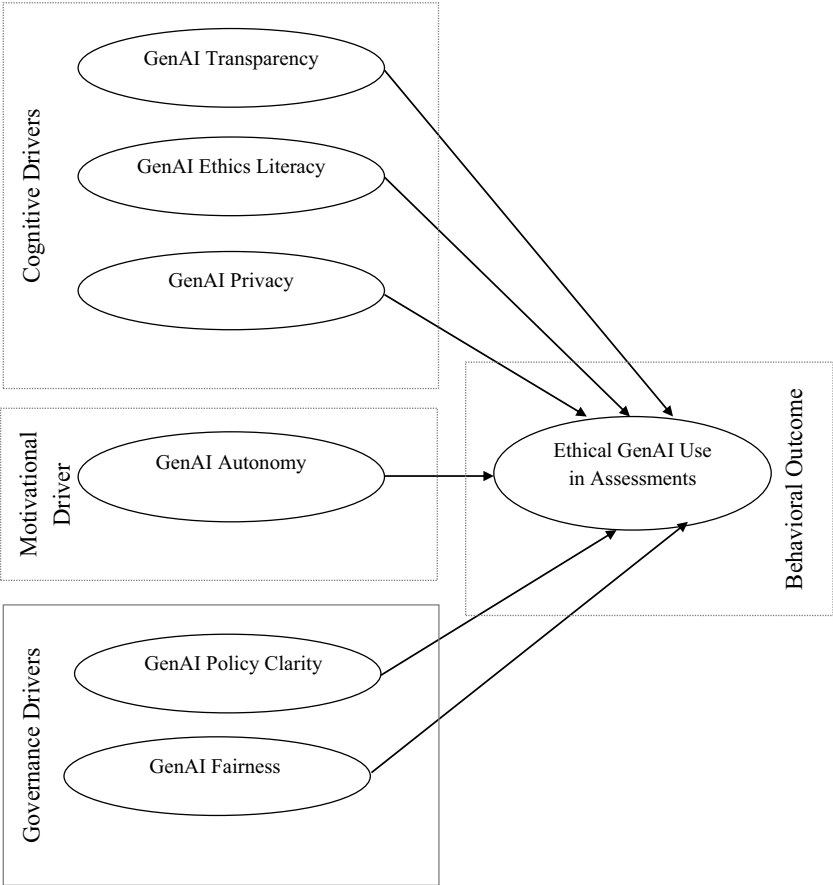


Figure 1. Model for ethical use of GenAI use in assessments

emerging domains such as ethical GenAI use, where constructs are relatively new. It can handle complex models with multiple latent constructs and paths, is robust with smaller sample sizes and non-normal data distributions and emphasizes prediction and variance explanation, which aligns with the study’s aim of explaining ethical behavior (Hair *et al.*, 2019; Hair and Alamer, 2022).

4.1 Measurement of items

This study used a structured questionnaire to collect demographic information from Malaysian university students. The questionnaire was developed around seven main constructs, comprising a total of 28 items, as outlined in Appendix. These constructs are: GenAI transparency, GenAI autonomy, GenAI privacy, GenAI ethics literacy, GenAI policy clarity, GenAI fairness and ethical GenAI use in assessments. The measurement scales were adapted from previously validated instruments (e.g. Shin, 2021; Chan, 2023; Huang, 2023; Collie and Martin, 2024; Baig and Yadegaridehkordi, 2025) and re-worded to align with the aims and scope of this study. The questionnaire items were measured using a five-point

Likert scale, where 1 = strongly disagree and 5 = strongly agree. To ensure clarity, relevance and appropriateness for the GenAI-based assessment context, the questionnaire items were reviewed by two field experts in educational assessment and GenAI. Their feedback was incorporated and the items were refined as needed to improve wording and grammatical accuracy. Meanwhile, the instrument was pilot-tested with a small group of participants before data collection and minor revisions were made to improve clarity and usability.

4.2 Data collection and sample size

This study focused on developing a theoretical framework for the ethical use of GenAI tools in academic assessments among university students. This study targeted students who have used GenAI because they are the primary users of these tools in academic assessments and can provide the most relevant experiential insights. The study focused on students from two leading Malaysian universities: Universiti Teknologi Malaysia and Universiti Malaya. Purposive sampling was used to ensure that respondents had direct exposure to academic integrity policies and GenAI use. Limiting participation to GenAI users also ensured that participants could meaningfully reflect on disclosure, attribution, acceptable assistance and boundary-setting in assessment contexts. While purposive sampling may limit generalizability, several steps were taken to mitigate bias:

- screening questions ensured only eligible participants proceeded; and
- participant anonymity and confidentiality protocols were implemented to minimize social desirability bias and encourage honest responses.

To verify eligibility, a screening question was placed at the beginning of the questionnaire, asking: “Have you ever used any types of GenAI in doing your assessment?” Only students who answered “yes” to this question were allowed to proceed to the full survey. The questionnaire was developed using Microsoft Forms and the link was distributed through university students’ social media platforms, including WhatsApp groups, to reach a broad student base. This approach was selected to maximize participation among the target population and is consistent with prior survey-based research in higher education ([Rabotapi and Matope, 2024](#)).

Respondents were informed that their participation was anonymous, that no personal identifying information was collected and that there were no right or wrong answers, thereby encouraging honest responses. A total of 171 complete and valid responses were collected. Following [Hair and Alamer’s \(2022\)](#) recommendation, the minimum sample size for PLS-SEM should be at least ten times the maximum number of structural paths pointing to any construct in the model. Therefore, the sample size of 171 exceeds the minimum requirement, ensuring the adequacy and robustness of the data set for PLS-SEM analysis ([Hair et al., 2019](#)).

4.3 Demographic data

According to the demographic profile ([Table 2](#)), the gender distribution consisted of 98 males (57.30%) and 73 females (42.69%). Regarding age, the respondents were primarily between 26 and 35 years (73 individuals, 42.69%), followed by 15–25 years (53 individuals, 30.99%), 36–50 years (34 individuals, 19.88%) and 50 years and above (11 individuals, 6.43%). In terms of educational attainment, a majority of respondents held undergraduate degrees (104 individuals, 60.81%), followed by graduate degrees (56 individuals, 32.74%), while a smaller portion reported other qualifications (11 individuals, 6.43%).

Regarding students’ first exposure to GenAI tools, 84 individuals (49.12%) had started using them within the past 7–12 months, 63 individuals (36.84%) had used them for over a year and 24 individuals (14.03%) had begun within the last 0–6 months. For self-reported

Table 2. Respondents demographics (*n* = 171)

Items	Categories	Frequency	%
Gender	Male	98	57.30
	Female	73	42.69
Age	15–25	53	30.99
	26–35	73	42.69
	36–50	34	19.88
	51 and above	11	6.43
	Undergraduate	104	60.81
Educational level	Graduate	56	32.74
	Other	11	6.43
When did students first start using GenAI?	0–6 months ago	24	14.03
	7–12 months ago	84	49.12
	More than a year ago	63	36.84
Self-reported proficiency level	Beginner	31	18.12
	Intermediate	76	44.44
	Advanced	64	37.42
Type of tools for assessment	ChatGPT	82	47.95
	Gemini	21	12.28
	Claude	16	9.35
	Bard	12	7.01
	Copy.ai	10	5.84
	All of the above	24	14.03
	Other	06	3.50
	Academic writing corrections	66	38.59
Purpose of use	Coding help	24	14.03
	Generating images	16	9.35
	Answering questions help	31	18.12
	Other	34	19.88

proficiency in using GenAI tools, 76 respondents (44.44%) identified as intermediate users, 64 respondents (37.42%) as advanced users and 31 respondents (18.12%) as beginners. In terms of the specific GenAI tools used for academic assessment purposes, ChatGPT was the most popular (82 individuals, 47.95%), followed by Gemini (21 individuals, 12.28%), Claude (16 individuals, 9.35%), Bard (12 individuals, 7.01%) and Copy.ai (10 individuals, 5.84%). In addition, 24 respondents (14.03%) reported using all listed tools, while six individuals (3.50%) used other tools. When asked about the purpose of using GenAI tools, the most common response was academic writing corrections (66 individuals, 38.59%), followed by answering questions help (31 individuals, 18.12%), coding help (24 individuals, 14.03%), generating images (16 individuals, 9.35%) and other purposes (34 individuals, 19.88%).

5. Results

In this study, the measurement model was assessed for reliability and validity. Reliability was evaluated using factor loadings (FL), Cronbach’s alpha (CA) and composite reliability (CR), while convergent validity was assessed through the average variance extracted (AVE). Discriminant validity was examined using the Fornell–Larcker criterion (FLC), cross-loadings (CL) and the heterotrait–monotrait (HTMT) ratio. The structural model was evaluated by examining collinearity using variance inflation factors (VIF), model fit using the standardized root mean square residual (SRMR), explanatory power (R^2), effect sizes

(f^2), predictive relevance (Q^2) and hypothesis testing. Together, these procedures provide a comprehensive validation strategy and demonstrate the model's reliability, validity, explanatory power and predictive relevance. This multi-step validation process enhances confidence in the robustness of the findings and addresses potential concerns regarding measurement bias and model misspecification (Hair *et al.*, 2019).

5.1 Measurement model

Hair and Alamer (2022) suggest that FL values should exceed 0.7. In this study, FL values ranged from 0.724 to 0.926, falling within the recommended range. Likewise, both CA and CR should exceed 0.7 to demonstrate reliability, which was confirmed as shown in Table 3. 'Convergent validity' was assessed by calculating AVE, with values above 0.5 for each variable indicating adequate convergent validity. Discriminant validity was assessed using the FLC, CL and the HTMT ratio.

The FLC was used to assess discriminant validity by comparing the AVE for each construct with its correlations with other constructs (see Table 4). According to this criterion, the square root of a construct's AVE should be greater than its correlations with all other constructs, indicating that the construct shares more variance with its own indicators than with other constructs in the model (Hair *et al.*, 2019). As shown in Table 4, the square roots

Table 3. Reliability and convergent validity

Constructs	Items	Factor loading (FL)	Cronbach's alpha (CA)	Composite reliability (CR)	Average variance extracted (AVE)
GenAI transparency	GT1	0.833	0.862	0.906	0.707
	GT2	0.803			
	GT3	0.852			
	GT4	0.875			
GenAI autonomy	GA1	0.913	0.905	0.932	0.774
	GA2	0.910			
	GA3	0.859			
	GA4	0.836			
GenAI privacy	GP1	0.860	0.898	0.929	0.765
	GP2	0.860			
	GP3	0.854			
	GP4	0.923			
GenAI ethics literacy	GEL1	0.926	0.907	0.935	0.783
	GEL2	0.826			
	GEL3	0.924			
	GEL4	0.861			
GenAI policy clarity	GPC1	0.902	0.880	0.916	0.732
	GPC2	0.880			
	GPC3	0.816			
	GPC4	0.822			
GenAI fairness	GF1	0.840	0.884	0.920	0.741
	GF2	0.856			
	GF3	0.873			
	GF4	0.874			
Ethical GenAI use in assessments	EGA1	0.856	0.802	0.870	0.627
	EGA2	0.773			
	EGA3	0.808			
	EGA4	0.724			

Table 4. Fornell-Larcker criterion

Constructs	EGA	GA	GEL	GF	GPC	GP	GT
1. Ethical GenAI use in assessments (EGA)	<i>0.792</i>						
2. GenAI autonomy (GA)	0.161	<i>0.880</i>					
3. GenAI ethics literacy (GEL)	0.509	0.347	<i>0.885</i>				
4. GenAI fairness (GF)	0.389	0.540	0.603	<i>0.861</i>			
5. GenAI policy clarity (GPC)	0.215	0.054	0.079	0.216	<i>0.856</i>		
6. GenAI privacy (GP)	0.384	0.290	0.433	0.275	0.004	<i>0.875</i>	
7. GenAI transparency (GT)	0.503	0.317	0.566	0.747	0.032	0.247	<i>0.841</i>

Note(s): Diagonal elements (in italics) represent the square roots of the average variance extracted

of the AVE values exceed the corresponding interconstruct correlations, confirming adequate discriminant validity.

CL was evaluated to determine whether each indicator had a higher loading on its intended construct compared to other constructs. According to the CL criterion, an indicator should not exhibit a higher loading on any construct other than its designated one (Hair and Alamer, 2022). As shown in Table 5, this criterion is satisfied in this study.

Table 5. Cross loadings

Indicators	EGA	GA	GEL	GF	GPC	GP	GT
EGA1	<i>0.856</i>	0.096	0.377	0.234	0.207	0.411	0.334
EGA2	<i>0.773</i>	0.180	0.549	0.526	0.512	0.218	0.628
EGA3	<i>0.808</i>	0.117	0.288	0.202	0.215	0.276	0.283
EGA4	<i>0.724</i>	0.102	0.353	0.200	0.095	0.332	0.268
GA1	0.140	<i>0.913</i>	0.303	0.518	0.076	0.283	0.311
GA2	0.184	<i>0.910</i>	0.344	0.467	0.060	0.329	0.283
GA3	0.114	<i>0.859</i>	0.216	0.417	0.010	0.180	0.209
GA4	0.104	<i>0.836</i>	0.351	0.514	0.053	0.135	0.315
GEL1	0.320	0.190	<i>0.926</i>	0.549	0.129	0.325	0.174
GEL2	0.319	0.280	<i>0.826</i>	0.554	0.074	0.145	0.143
GEL3	0.327	0.264	<i>0.924</i>	0.527	0.103	0.325	0.322
GEL4	0.373	0.277	<i>0.861</i>	0.512	0.369	0.639	0.225
GF1	0.418	0.329	0.514	<i>0.840</i>	0.373	0.478	0.482
GF2	0.400	0.367	0.471	<i>0.856</i>	0.555	0.365	0.474
GF3	0.514	0.277	0.690	<i>0.873</i>	0.652	0.471	0.541
GF4	0.456	0.273	0.284	<i>0.874</i>	0.419	0.470	0.499
GPC1	0.171	0.032	0.369	0.683	<i>0.902</i>	0.569	0.013
GPC2	0.158	0.104	0.014	0.487	<i>0.880</i>	0.320	0.011
GPC3	0.153	0.025	0.016	0.742	<i>0.816</i>	0.025	0.056
GPC4	0.230	0.031	0.489	0.577	<i>0.822</i>	0.395	0.646
GP1	0.270	0.453	0.541	0.518	0.351	<i>0.860</i>	0.730
GP2	0.346	0.413	0.641	0.722	0.425	<i>0.860</i>	0.581
GP3	0.382	0.520	0.351	0.577	0.325	<i>0.854</i>	0.622
GP4	0.325	0.469	0.265	0.508	0.147	<i>0.923</i>	0.526
GT1	0.433	0.323	0.142	0.710	0.258	0.512	<i>0.833</i>
GT2	0.406	0.161	0.325	0.526	0.453	0.426	<i>0.803</i>
GT3	0.436	0.352	0.125	0.314	0.480	0.324	<i>0.852</i>
GT4	0.416	0.219	0.140	0.096	0.413	0.410	<i>0.875</i>

Note(s): Ethical GenAI use in assessments (EGA), GenAI autonomy (GA), GenAI ethics literacy (GEL), GenAI fairness (GF), GenAI policy clarity (GPC), GenAI privacy (GP), GenAI transparency (GT)

As shown in Table 6, all HTMT values are below 1.0, providing further evidence of discriminant validity. According to Henseler *et al.* (2015), HTMT values close to 1.0 suggest insufficient discriminant validity, whereas values substantially below 1.0 indicate that the constructs are empirically distinct.

5.2 Structural model

5.2.1 Collinearity assessment. The first step in evaluating the structural model is to assess potential collinearity among the predictor constructs. VIF values were examined to identify multicollinearity, as VIF indicates the degree to which variance in a predictor is inflated by its correlation with other predictors (Hair and Alamer, 2022). According to Joseph *et al.* (2017), tolerance values greater than 0.20 and VIF values below 5.0 are considered acceptable, whereas higher VIF values may indicate problematic collinearity that could compromise the stability of the path estimates. Table 7 shows that VIF values ranged from 1.11 (GenAI Policy Clarity) to 3.64 (GenAI Fairness), all well below the recommended threshold of 5. This indicates that multicollinearity is not a concern; consequently, no predictor constructs require removal or adjustment.

5.2.2 Model fit. The overall model fit was assessed using the SRMR, which measures the average magnitude of differences between the observed correlations and those predicted by the structural model. SRMR is widely used in PLS-SEM to evaluate the degree to which the model reproduces the empirical correlation matrix (Joseph *et al.*, 2017). According to Hu and Bentler, (1999), SRMR values below 0.08 indicate good model fit, whereas values above 0.08 suggest poor fit. In the present study, the SRMR value was 0.07, which is below the recommended threshold. This result indicates that the model exhibits a good fit, demonstrating that the relationships specified in the structural model adequately represent the observed data.

5.2.3 Hypotheses testing. Following the 95% confidence level guideline by Hair *et al.* (2019), the structural model was assessed to evaluate the proposed hypotheses. Table 8 and

Table 6. Heterotrait-monotrait ratio (HTMT)

Constructs	1	2	3	4	5	6	7
1. Ethical GenAI use in Assessments							
2. GenAI autonomy	0.175						
3. GenAI ethics literacy	0.575	0.385					
4. GenAI fairness	0.429	0.605	0.678				
5. GenAI policy clarity	0.244	0.071	0.109	0.240			
6. GenAI privacy	0.459	0.303	0.476	0.309	0.044		
7. GenAI transparency	0.574	0.356	0.637	0.621	0.079	0.282	

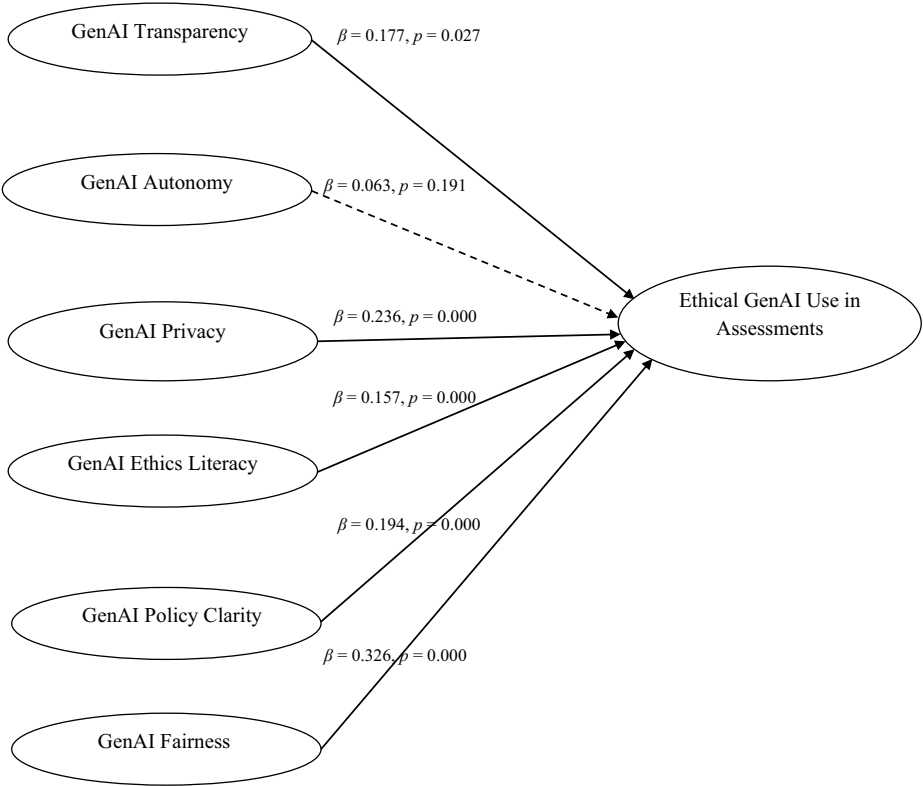
Table 7. Collinearity statistics measurement

Constructs	VIF
GenAI autonomy	1.52
GenAI ethics literacy	1.97
GenAI fairness	3.64
GenAI policy clarity	1.11
GenAI privacy	1.27
GenAI transparency	3.20

Table 8. Hypotheses testing

Hypotheses	β	p -values	Status	Decision
H1. GenAI transparency → Ethical GenAI use in assessments	0.177	0.027	$p < 0.05$	Accepted
H2. GenAI autonomy → Ethical GenAI use in assessments	0.063	0.191	$p > 0.05$	Rejected
H3. GenAI privacy → Ethical GenAI use in assessments	0.236	0.000	$p < 0.05$	Accepted
H4. GenAI ethics literacy → Ethical GenAI use in assessments	0.157	0.000	$p < 0.05$	Accepted
H5. GenAI policy clarity → Ethical GenAI use in assessments	0.194	0.000	$p < 0.05$	Accepted
H6. GenAI fairness → Ethical GenAI use in assessments	0.326	0.000	$p < 0.05$	Accepted

Figure 2 summarize the hypothesis testing outcomes, detailing the path coefficients (β) and corresponding p-values. Significant relationships were found for H1 ($\beta = 0.177, p = 0.027$), H3 ($\beta = 0.236, p = 0.000$), H4 ($\beta = 0.157, p = 0.000$), H5 ($\beta = 0.194, p = 0.000$) and H6 ($\beta = 0.326, p = 0.000$), leading to the acceptance of these hypotheses. In contrast, H2 ($\beta = 0.063, p = 0.191$) exhibited an insignificant relationship and was therefore rejected



Note: Dashed arrows indicate rejected hypotheses

Figure 2. Hypotheses results

Figure 2. In addition, the R^2 value was used to evaluate the overall strength and explanatory power of the model. In this study, the R^2 value for this model is 0.556. According to Chin (1998), an R^2 exceeding 0.33 indicates a significant model. Thus, the R^2 values for the proposed model are deemed significant.

5.2.4 Effect size (f^2). The effect size (f^2) was computed to measure the impact of each predictor on the endogenous constructs (Cohen, 1988; Hair *et al.*, 2017). According to Cohen, f^2 values of 0.02, 0.15, and 0.35 represent small, medium and large effects, respectively. As presented in Table 9, GenAI privacy ($f^2 = 0.99$), GenAI policy clarity ($f^2 = 0.76$) and GenAI fairness ($f^2 = 0.66$) exhibit large effect sizes, indicating that these constructs are the most influential predictors of ethical GenAI use in assessments. In contrast, GenAI ethics literacy ($f^2 = 0.28$), GenAI autonomy ($f^2 = 0.26$) and GenAI transparency ($f^2 = 0.22$) demonstrate moderate effect sizes, reflecting comparatively lower but still meaningful contributions.

Although the structural path of autonomy was statistically insignificant, the predictor demonstrated a moderate effect size ($f^2 > 0.15$). Prior PLS-SEM literature clarifies that statistical significance and effect size assess different aspects of structural relationships and therefore may yield differing results (Chin, 1998; Hair *et al.*, 2021). Specifically, while path significance assesses whether a relationship is statistically detectable, effect size (f^2) evaluates the relative contribution of an exogenous construct to the explained variance of the endogenous construct. Consequently, a predictor may exert a practically important influence on the model despite an insignificant path coefficient, particularly in exploratory or behavioral research contexts.

5.2.5 Predictive relevance Q^2 . The predictive relevance of the model was assessed using the Q^2 statistic obtained through the cross-validated redundancy approach, as recommended by Hair *et al.* (2019). According to Hair and Alamer (2022), values of Q^2 greater than zero indicate that a model has predictive relevance. The Q^2 value was calculated using the blindfolding procedure in SmartPLS. The results show that “ethical GenAI use in assessments” achieved a Q^2 value of 0.308, demonstrating that the model has substantial predictive relevance for the construct (Table 10).

As reported in Tables 9 and 10, the f^2 values indicate medium and large effects, whereas Q^2 value is greater than zero, suggesting substantial explanatory power and strong contributions to predictive relevance.

Table 9. Effect size (f^2)

Constructs	f^2 (effect size)	Category
GenAI privacy	0.99	Large
GenAI policy clarity	0.76	Large
GenAI fairness	0.66	Large
GenAI ethics literacy	0.28	Medium
GenAI autonomy	0.26	Medium
GenAI transparency	0.22	Medium

Table 10. Predictive relevance Q^2 results

Construct	Predictive relevance Q^2
Ethical GenAI use in assessments	0.308

6. Discussions

This study aimed to investigate the students' ethical use of GenAI tools in educational assessments, focusing on key factors such as transparency, autonomy, privacy, ethics literacy, policy clarity and fairness. The analysis showed that privacy emerged as the strongest predictor of ethical GenAI use, demonstrating the largest effect size among all examined factors. This finding aligns with both the normative ethics framework (Huang, 2023) and digital literacy theory, suggesting that students who understand privacy risks are more likely to engage with GenAI responsibly. Ethical integration of GenAI in assessment contexts depends not only on system functionality but also on the safeguarding of students' personal and academic data. When privacy protection is weak, risks such as data breaches, unauthorized surveillance or misuse of assessment-related information may arise, potentially undermining trust and encouraging disengagement or unethical practices. The findings align with prior research emphasizing that robust data governance and privacy assurance are foundational conditions for ethical AI use in educational settings (Huang, 2023; Marín *et al.*, 2025).

GenAI policy clarity also showed a significant and large effect on the ethical use of GenAI in assessments. This finding highlights the importance of clear institutional guidelines in shaping responsible behavior, aligning with SDT, which emphasizes the role of structured environments in supporting effective decision-making. Well-detailed policies serve as a reference point for students, helping them distinguish between appropriate support and academic misconduct when using AI tools. Without explicit guidance, students may unknowingly misuse GenAI, leading to ethical breaches such as plagiarism or over-reliance on AI-generated content. Policy clarity not only reduces ambiguity but also fosters a culture of accountability and informed decision-making among students. This finding is consistent with Chan (2023), who argues that the absence of clear institutional norms often leads to inconsistent practices and undermines trust in AI-supported education systems.

Fairness further demonstrated a large effect on ethical GenAI use by ensuring that AI systems operate equitably and without bias. This finding is supported by the normative ethics framework (Huang, 2023), which emphasizes fairness as a core principle for responsible AI use. It also aligns with digital literacy theory, as students with higher literacy are better able to recognize potential biases and inequities in AI outputs. Together, these perspectives suggest that fairness plays a critical role in guiding ethical decision-making and promoting responsible engagement with GenAI tools. This result reinforces the view that fairness is a foundational principle in the ethical application of AI, especially in academic assessment contexts where equitable treatment is essential. When GenAI tools operate without bias, they help ensure that no student is unfairly advantaged or disadvantaged based on factors such as language proficiency, cultural background or learning style. This finding aligns with research on fairness in educational AI, which emphasizes the need to mitigate bias to ensure ethical integrity and equal treatment for all students (Memarian and Doleck, 2023).

GenAI ethics literacy contributed significantly to responsible AI use, demonstrating a moderate effect size while serving as an important cognitive safeguard. This finding is consistent with digital literacy theory, which posits that ethics literacy enhances students' ability to recognize, evaluate and respond to ethical issues associated with AI use. Ethics literacy equips students with knowledge and evaluative skills needed to navigate complex ethical dilemmas in assessments (Foltynek *et al.*, 2023). Although its effect is moderate, ethics literacy interacts with stronger structural predictors such as policy clarity, fairness and privacy, enabling fully informed and responsible use. This finding aligns with previous studies highlighting the critical role of ethics training in shaping student behavior and moral judgment in technology-mediated learning contexts (Bego *et al.*, 2024).

The analysis showed that GenAI transparency plays an important role in promoting the ethical use of AI tools in academic assessments, with a medium effect size indicating a meaningful and practically significant influence. Consistent with digital literacy theory, transparency enhances students' ability to understand and critically evaluate AI-generated outputs, thereby supporting informed and responsible decision-making. When students understand how AI systems generate content and are aware of their limitations, they are better equipped to critically evaluate AI-generated information and integrate AI assistance appropriately. Transparency acts as a foundational cognitive safeguard, helping reduce misinterpretation, over-reliance and potential misuse of AI-generated material. This is consistent with prior research identifying transparency as a key factor supporting informed decision-making and responsible AI engagement in educational contexts (Choi *et al.*, 2024; Lazăr *et al.*, 2024).

Finally, autonomy, representing motivational capacity, showed a nonsignificant but moderate effect on the ethical use of GenAI in assessment. While GenAI autonomy was conceptualized as learners' capacity to regulate their use of these tools responsibly within institutional constraints, this finding suggests that autonomy alone may be insufficient to ensure ethical engagement. From an SDT perspective, autonomy supports self-regulated behavior only when accompanied by sufficient competence and appropriate structure. In contexts where students lack clear guidelines, ethical awareness or the necessary digital literacy to critically evaluate AI outputs, autonomy may not translate into responsible behavior. Instead, it may lead to inconsistent or suboptimal decision-making. This may explain why students who feel empowered to use GenAI do not necessarily engage in careful judgment, critical evaluation or adherence to academic integrity standards. In the absence of complementary supports such as clear assessment policies, explicit expectations and ethical literacy, autonomy may facilitate convenience-driven behaviors, including over-reliance on AI-generated content or superficial engagement. Consistent with He and Li (2023), this finding highlights that while autonomy can enhance flexibility and efficiency, it must be balanced with guidance and constraints to foster responsible and accountable GenAI use in academic assessments. Therefore, the nonsignificant effect observed in this study may reflect a misalignment between perceived autonomy and the structural and cognitive conditions required for ethical decision-making.

In summary, students' ethical use of GenAI in assessments is primarily driven by competency-based factors (from digital literacy), including fairness, privacy, policy clarity, transparency and ethics literacy, while perceived autonomy (from SDT) alone is not significant. This highlights that ethical engagement arises from the interaction of knowledge, skills and contextual support, while autonomy requires complementary structures to be effective. Emphasizing the relative strengths of predictors provides a nuanced understanding of how theoretical constructs jointly shape ethical decision-making in AI-mediated assessments.

7. Implications

7.1 Theoretical implications

This study contributes to the growing body of research on ethical GenAI use in assessment by developing and empirically testing an integrated theoretical framework grounded in digital literacy theory and SDT. Specifically, the study identifies key determinants of ethical GenAI behavior, including GenAI ethics literacy, transparency, privacy, fairness, policy clarity and autonomy, with these constructs further informed by Huang's (2023) normative AI ethics framework. In doing so, the study translates abstract ethical principles into empirically measurable factors.

By integrating digital literacy theory and SDT, the study offers a complementary perspective that captures both competence-based and motivational explanations of ethical behavior. Digital literacy theory explains how students' knowledge and skills, particularly GenAI ethics literacy, shape their ability to critically evaluate and engage with AI technologies. In parallel, SDT provides a motivational lens, highlighting the role of autonomy-supportive conditions in facilitating the internalization of ethical values within GenAI-enabled assessment contexts.

Through this theoretical integration, the study moves beyond prior GenAI assessment research, which has largely been descriptive or exploratory, by offering a structured framework for understanding students' ethical decision-making. The study conceptualizes ethical GenAI use as an outcome of the interplay among competence, motivational internalization and ethical governance mechanisms. This perspective advances the framework from a simple integration of variables to a more substantive behavioral explanation of ethical AI use in assessment contexts. The empirical findings support most of the proposed relationships; however, the effect of GenAI autonomy on ethical GenAI use was not significant. This result provides an important theoretical insight by refining the application of SDT in technology-mediated assessment contexts. While SDT emphasizes the central role of autonomy in fostering internalization and self-regulated behavior, the findings suggest that the influence of autonomy may be context dependent.

Specifically, in structured assessment environments characterized by clear rules, accountability and evaluative consequences, the direct effect of autonomy on ethical behavior appears to be reduced. Instead, governance-related factors such as policy clarity, fairness, privacy and transparency, along with students' ethical competence, play a more prominent role in shaping behavior. This implies that autonomy may operate indirectly, potentially in interaction with competence-related factors, rather than functioning as an independent predictor of ethical GenAI use. This insight contributes to theory by demonstrating that motivational constructs such as autonomy must be considered alongside governance and contextual constraints when explaining ethical AI behavior in assessment settings.

Overall, the study advances theory by demonstrating that ethical GenAI use is best understood through the integration of cognitive (literacy-based), motivational (autonomy-related) and governance-related (transparency, fairness, privacy and policy clarity) factors. By examining these determinants within a unified framework, the study provides a more comprehensive and empirically grounded understanding of ethical AI behavior in higher education assessment contexts.

7.2 Practical implications

This study presents several key practical implications for students, institutions, educators and GenAI developers, based on the hypotheses tested concerning ethical GenAI use in assessments. This study shows that transparency plays a crucial role in promoting the ethical use of GenAI in academic assessments, offering several important practical contributions for students, educators, institutions and GenAI developers. When GenAI use is transparent, students can approach AI-supported work with deliberate awareness of how outputs are generated. For example, understanding that GenAI relies on predictive text generation rather than verified reasoning enables students to examine outputs more carefully before incorporating them into academic work. Transparency also motivates students to verify AI-generated elements, such as checking whether references actually exist, confirming the accuracy of theoretical explanations and validating empirical claims, thereby reducing the risk of submitting fabricated or misleading information (Kim *et al.*, 2025). In addition, students can routinely disclose GenAI assistance in their assignments by including a short

GenAI use statement that specifies which parts of the task were supported by AI (e.g. brainstorming, proofreading, structuring). Institutions can use these insights to strengthen academic integrity frameworks by implementing policies requiring GenAI transparency statements in all assessments, providing institutional templates for AI-use disclosures and aligning ethical AI guidelines (Lim *et al.*, 2025). Institutions can also integrate training modules into orientation programs to ensure consistency in practice across departments. Educators can enhance assessment design by embedding transparency requirements, such as requiring students to document how GenAI contributed to their drafts, justifying the validity of AI-generated content or reflecting on the tool's limitations. They can also guide students in critically evaluating AI outputs by modeling verification techniques during class activities. Finally, GenAI developers can support ethical academic use by improving transparency and explainability features. This includes publishing accessible documentation on model training data, known limitations and typical error patterns; integrating features that explain outputs and allow users to view reasoning summaries; and implementing notifications that encourage proper disclosure in academic settings. The effectiveness of transparency mechanisms can be assessed using indicators such as AI-use disclosure accuracy, reporting compliance rates and reduction in undetected AI-generated content in submitted assessments.

The structural model demonstrates that privacy perceptions significantly influence ethical usage behavior, suggesting that strengthened data protection mechanisms and clear communication about data handling can reinforce trust-based compliance among students. Accordingly, institutional data protection frameworks function not only as compliance requirements but also as behavioral governance mechanisms that sustain trust-based ethical engagement in assessment contexts. Students should understand how their data are handled and follow institutional guidelines for responsible AI use. They can adopt practical measures, such as avoiding the sharing of sensitive personal information in AI prompts and reviewing AI outputs before submission (Huang, 2023; Lim *et al.*, 2025). Institutions should implement secure submission platforms and monitor compliance with privacy protocols and ensure adherence to recognized data protection standards. They can establish auditing mechanisms to review AI tools' data practices and enforce accountability. By codeveloping privacy policies with faculty, institutions can clearly outline permissible AI usage, data storage practices and student rights, thereby fostering trust and promoting ethical use of GenAI technologies. For instance, annual audits of AI platforms and student-facing portals explaining data protection measures can provide transparency and confidence. Educators should engage in professional development programs that address AI data privacy and ethical assessment practices. They can integrate privacy considerations into assignment design by prompting students to anonymize sensitive inputs and reflect on the AI's contribution without disclosing personal and institutional information. Clear communication about how student data is collected, stored and used reinforces trust and encourages responsible engagement (Huang *et al.*, 2024). GenAI developers should embed privacy-by-design principles into AI tools, including anonymization of student inputs, minimization of data retention and secure storage. Comprehensive documentation detailing what data is collected, how it is processed and retention policies can support ethical use in educational contexts. Developers can also provide guidance for secure and responsible use in assessments and create features that allow institutions, educators and students to monitor privacy compliance. For example, a privacy dashboard within the AI tool can show what data is stored, for how long and provide options for deletion or export, fostering trust and ethical engagement. Privacy effectiveness can be evaluated through measurable indicators such as data-sharing incidents, privacy violation rates and compliance with institutional data protection standards.

GenAI ethics literacy is essential for enabling students to engage responsibly with AI tools in academic assessments. Students who possess stronger ethical awareness and competence are better equipped to recognize risks such as bias, fairness concerns and accountability issues, all of which influence how they interact with AI technologies (Wiese *et al.*, 2025). By developing ethics literacy, students learn to critically evaluate AI-generated outputs, for example, by checking whether citations are authentic, identifying potential bias in explanations or detecting oversimplified reasoning patterns. This allows them to reflect on the ethical implications of using AI and apply principles of responsibility, inclusivity and academic integrity to ensure that AI-assisted work remains fair and trustworthy. To foster this capability, educational institutions should invest in structured training and professional development programs that focus on AI ethics and responsible use. Examples include offering AI ethics workshops for new students, embedding ethics modules into core units or requiring students to complete short training quizzes on bias, transparency and responsible disclosure before submitting assessments. Educators play a pivotal role in developing students' ethics literacy. They can be trained to help students identify ethical challenges in AI outputs, such as biased language or fabricated references and integrate ethical evaluation into assignments. Practical strategies include showing students examples of incorrect or biased AI-generated content, modeling how to verify or correct such outputs and incorporating reflective prompts into assignments that allow students to explain how they validated the accuracy and fairness of the AI-generated material used in their assessments. By modeling responsible practices and emphasizing ethical reasoning in feedback, educators help students build the skills needed to engage critically with GenAI. Developers can also support ethical engagement by incorporating user-oriented ethical guidance directly into AI tools. Examples include tooltips that remind users to verify references, prompts that encourage reflection on potential bias or built-in indicators that highlight sections requiring further human validation. Such features reinforce ethical awareness at the point of use and help promote responsible academic application of GenAI technologies. Ethics literacy outcomes can be measured using assessment scores in AI ethics training, quality of reflective responses and students' ability to identify bias or inaccuracies in AI outputs.

Clear and well-defined policies established by educational institutions help ensure the consistent, fair and transparent use of GenAI across academic contexts. The empirical effect of policy clarity on ethical GenAI usage suggests that reducing ambiguity through explicit definitions of permissible and impermissible AI use directly stabilizes student decision-making and decreases the likelihood of policy violations (An *et al.*, 2025). Embedding these definitions within course syllabi and assessment guidelines therefore serves as an important structural mechanism for reducing interpretive uncertainty surrounding GenAI use. When policies clearly outline expectations for AI-assisted work, students can better understand the boundaries of acceptable behavior, align their actions with ethical and academic standards and avoid unintentional misuse. Effective guidelines articulate ethical responsibilities, specify appropriate disclosure requirements and outline consequences for improper use, supporting consistent application across programs and learning environments (Chan, 2023; Dabis and Csáki, 2024). Institutions can strengthen ethical and responsible GenAI use by establishing comprehensive policies that define permissible forms of AI assistance, expectations for transparency and standards for academic integrity in AI-supported assessments. Educators benefit from professional development that equips them to interpret and implement these policies, design assessments aligned with institutional expectations, guide students in ethical AI engagement and provide feedback that reinforces policy adherence. For developers, GenAI systems can be designed to support institutions in operationalizing assessment rules through configurable usage settings, clearly defined

permissible output modes, monitoring capabilities and reporting mechanisms. Such features help ensure that GenAI tools align with institutional integrity frameworks and promote responsible academic use. Policy effectiveness can be evaluated using compliance rates, frequency of violations and consistency of AI-use disclosures across academic programs.

GenAI fairness is a fundamental element in promoting ethical engagement with AI tools in academic assessments, as perceptions of equity strongly influence responsible use and trust. The significant association between fairness perceptions and ethical behavior suggests that institutional monitoring of bias and ensuring equitable access to GenAI tools can serve as structural mechanisms that reinforce trust-based ethical engagement. When students perceive GenAI systems as fair and capable of producing unbiased outcomes, their confidence in assessment processes increases, which in turn encourages responsible and transparent use (Memarian and Doleck, 2023). Students can be guided to critically evaluate AI-generated outputs and reflect on whether the AI-assisted components of their work uphold fairness and inclusivity. Institutions can reinforce fairness by adopting GenAI tools that incorporate built-in bias mitigation features. Policies and oversight procedures can explicitly address fairness by ensuring that all students, regardless of discipline, ability level or background, have equal access to AI tools and receive consistent guidance on their ethical use (An *et al.*, 2025). An example of this is providing institution-wide access to approved GenAI platforms rather than allowing only some students to use paid or external tools. Educators play a complementary role by guiding students to critically appraise AI outputs, integrating fairness considerations into assessment design and reinforcing equitable evaluation practices within the classroom. GenAI developers also contribute by embedding bias mitigation into system design, conducting regular audits and providing transparent documentation of fairness measures and limitations, ensuring that the tools themselves operate equitably and align with institutional and classroom expectations. Fairness outcomes can be assessed through student perception surveys, equity of access indicators and bias audit results of AI systems.

Generally, detection-based approaches to academic integrity are widely recognized as costly, adversarial and increasingly limited in effectiveness as AI tools evolve (Dawson, 2020; Corbin *et al.*, 2025). Collectively, the identified predictors suggest that preventive governance strategies grounded in literacy development, transparency, fairness and policy clarity may be more effective in shaping ethical behavior than reactive surveillance-based enforcement mechanisms. From a cost–benefit perspective, the proposed model supports institutions in prioritizing investments in governance, ethics literacy and assessment redesign that align with university policies, rather than relying on reactive enforcement mechanisms. By embedding ethical principles at the design stage and strengthening students' skills and knowledge related to ethical GenAI use, universities may reduce long-term costs associated with academic misconduct, policy violations, reputational damage and GenAI tool misuse.

Taken together, these recommendations can be understood as interconnected components of an ethical GenAI governance framework, illustrating how institutions can translate theoretical constructs into actionable interventions while ensuring consistency, accountability and fairness in assessment contexts.

7.2.1 Social implications. Universities play a critical role in preparing students for participation in an AI-driven society (Magrill and Magrill, 2024). Therefore, at a broader societal level, the findings of this study suggest that fostering ethical GenAI use among students contributes to the development of digitally responsible graduates who are better prepared to engage with AI technologies beyond higher education. As GenAI tools become embedded in professional and everyday decision-making, early exposure to transparent, fair and privacy-aware AI practices may shape public attitudes toward responsible AI use, supporting trust without uncritical dependence (Huang, 2023). When students perceive GenAI use as

transparent, fair and governed by clear policies, they are more likely to internalize ethical norms, as evidenced by the significant effects of transparency, fairness and policy clarity in the model, thereby strengthening trust in assessment outcomes and academic qualifications. Meanwhile, by emphasizing fairness, transparency and ethics literacy, universities can help reduce social disparities and ensure that GenAI-enhanced assessment does not advantage particular student groups. This implication is supported by the statistically significant relationship between fairness and ethical GenAI use, linking equity concerns to observed behavioral outcomes.

From a policy perspective, these findings support the development of structured GenAI governance frameworks that go beyond general guidelines to include enforceable assessment standards, standardized AI-use disclosure mechanisms and institutional monitoring systems. The model confirmed that policy clarity significantly reduces ambiguity in student decision-making, as demonstrated by its significant predictive effect on ethical GenAI behavior, providing empirical support for enforceable disclosure mechanisms and monitoring systems. These can be implemented through learning management systems that track AI-assisted submissions and generate compliance reports. Policy effectiveness can be evaluated using indicators such as disclosure rates, policy violation frequency and adherence to institutional GenAI guidelines. From an economic perspective, preventive governance strategies based on ethics literacy, transparency and structured assessment design can reduce costs associated with misconduct investigations, manual verification and enforcement of AI-related violations. This is substantiated by the positive effect of ethics literacy on ethical GenAI use, showing that preventive strategies grounded in literacy development directly reduce misconduct risks.

At a societal level, embedding ethical GenAI practices in higher education contributes to preparing graduates for AI-integrated environments. Students exposed to structured ethical AI use are more likely to demonstrate responsible digital behavior in professional contexts, which strengthens public trust in AI systems. This aligns with the model's findings that ethically guided conditions (e.g. transparency and fairness) significantly influence behavior, reinforcing the societal relevance of these constructs. In addition, equitable access to GenAI tools and consistent ethical guidance can reduce digital inequality and promote inclusive participation in AI-enabled learning and work environments. The empirical association between fairness perceptions and ethical behavior reinforces this point, evidencing that fairness is not only a normative principle but a statistically supported driver of responsible engagement. Collectively, these outcomes enhance societal readiness for responsible AI adoption and support sustained trust in AI-driven systems across education and industry.

8. Conclusion

This study developed and empirically tested a theoretical framework supporting students' ethical use of GenAI tools in academic assessments, integrating digital literacy theory and SDT. Key findings reveal that GenAI transparency, privacy, ethics literacy, policy clarity and fairness significantly promote ethical engagement, with relative strengths highlighting their centrality in guiding responsible AI use. In contrast, GenAI autonomy did not have a direct effect. This suggests that independent decision-making alone is insufficient and should be supported by ethics literacy and institutional governance mechanisms to foster ethical GenAI use.

These results provide actionable insights for stakeholders who embed ethical principles into assessment design and support students' digital literacy to foster predictable, responsible GenAI behavior systematically. The study also offers a theoretically grounded, explanatory model linking factors to ethical AI engagement in academic contexts.

9. Limitations and future directions

This study provides valuable insights into a theoretical framework for supporting students' ethical use of GenAI tools in academic assessments. However, certain limitations must be acknowledged. The sample was limited to students with prior GenAI experience from two Malaysian universities, selected through purposive sampling, which may constrain the generalizability of the findings. Importantly, the findings are context-specific to Malaysian higher education and should not be interpreted as universally generalizable. The institutional policies, regulatory environment, digital infrastructure and cultural norms surrounding academic integrity in Malaysia may differ substantially from those in other countries. Therefore, the proposed framework does not have universal international applicability. Replication with nonusers, additional institutions and broader national or international samples is needed to assess the generalizability of the proposed model further, as findings may vary across contexts due to differences in policy clarity, digital infrastructure and cultural perceptions of academic integrity. The reliance on purposive sampling of GenAI users may also introduce selection bias, as students who already use GenAI tools may demonstrate higher levels of digital literacy, ethical awareness or technology acceptance than nonusers. Consequently, the relationships observed in this study may not fully represent the broader student population. Future research should incorporate more diverse and probabilistic sampling strategies, including both users and nonusers, to examine potential differences in ethical orientations and behaviors. In addition, the proposed model, grounded in digital literacy and SDT, focuses on five ethical factors (transparency, privacy, ethical awareness, policy clarity and fairness). While these constructs capture key ethical considerations relevant to student decision-making, they may not cover the full assessment ecosystem. For example, in this study, only the autonomy component of SDT was used as the most direct antecedent of ethical GenAI use in assessments. Future research could extend the model by incorporating competence- and relatedness-related measures to examine additional motivational pathways underlying ethical GenAI use. In addition, developing a multi-level model incorporating mediation analysis would allow researchers to capture cross-level influences and examine more complex explanatory mechanisms underlying ethical GenAI use. Future studies can also extend and validate the framework by incorporating additional ethical dimensions and stakeholder perspectives (e.g. instructors and administrators) across diverse institutions. This study used a quantitative approach to examine relationships among key constructs. Future research may incorporate qualitative methods to explore how students interpret and negotiate ethical GenAI use, as well as to examine the content and implementation of policy clarity and ethics literacy initiatives. Future studies could extend this work by using more advanced analytical techniques to deepen the understanding of the proposed relationships. For example, multi-group analysis can be used to examine whether the effects of GenAI transparency differ across demographic groups, academic disciplines or levels of GenAI experience. Researchers may also consider assessing predictive validity by testing whether the constructs in this model reliably forecast students' future ethical behaviors or long-term patterns of GenAI use. This study used indirect indicators to approximate students' actual ethical behaviors because direct behavioral observation was not feasible within the scope of the research design. Future studies could directly measure real student behaviors by capturing how students actually interact with GenAI platforms, analyzing authentic usage patterns and monitoring disclosure practices. Another limitation relates to the potential influence of social desirability bias and common method bias. As the study relies on self-reported survey data, respondents may have been inclined to provide socially acceptable answers, particularly given the ethical nature of GenAI use in assessments. This may have led to an overestimation of ethical behavior. In addition,

collecting data on both predictors and outcomes from the same source at a single point in time raises concerns regarding common method variance. Although procedural remedies, such as assuring respondent anonymity and carefully designing survey items, were implemented to mitigate these effects, they cannot be eliminated. Future research should consider using multi-source data, behavioral measures or longitudinal designs to further reduce the risk of bias and strengthen the robustness of the findings. Combining multiple data sources will enhance the accuracy and robustness of behavioral measurement in future work. Finally, this study focused on the ethical use of GenAI in assessments at a particular point in time. As new GenAI tools are developing rapidly, longitudinal studies could track shifts in attitudes, policy interpretations and ethical concerns over time, providing insights into the long-term effectiveness and challenges of implementing GenAI in assessment contexts.

References

- Adarkwah, M.A. (2025), "GenAI-Infused adult learning in the digital era: a conceptual framework for higher education", *Adult Learning*, Vol. 36 No. 3, pp. 149-161.
- Agustini, N.P.O. (2023), "Examining the role of ChatGPT as a learning tool in promoting students' English language learning autonomy relevant to Kurikulum Merdeka Belajar", *Edukasia: Jurnal Pendidikan Dan Pembelajaran*, Vol. 4 No. 2, pp. 921-934.
- Al-Kfairy, M., Mustafa, D., Kshetri, N., Insiew, M. and Alfandi, O. (2024), "Ethical challenges and solutions of generative AI: an interdisciplinary perspective", *In Informatics*, Vol. 11 No. 3, p. 58. Multidisciplinary Digital Publishing Institute.
- An, Y., Yu, J.H. and James, S. (2025), "Investigating the higher education institutions' guidelines and policies regarding the use of generative AI in teaching, learning, research, and administration", *International Journal of Educational Technology in Higher Education*, Vol. 22 No. 1, p. 10.
- Baig, M.I. and Yadegaridehkordi, E. (2025), "Factors influencing academic staff satisfaction and continuous usage of generative artificial intelligence (GenAI) in higher education", *International Journal of Educational Technology in Higher Education*, Vol. 22 No. 1, p. 5.
- Bego, C.R., Withorn, T., Danovitch, J., Thompson, A., Thomas, E., Gatsos, G.E. and Tran, A. (2024), "Working towards GenAI literacy: assessing first-year engineering students' attitudes towards, trust in, and ethical opinions of ChatGPT", *2024 ASEE annual conference and exposition*.
- Belshaw, D. (2012), "'What is' digital literacy'? A pragmatic investigation", Doctoral dissertation, Durham University.
- Beynen, T. (2024), "The role of students' assessment literacies in navigating university assessment, GenAI, and academic integrity", *Brock Education: A Journal of Educational Research and Practice*, Vol. 33 No. 3, pp. 30-56.
- Bower, M., Torrington, J., Lai, J.W., Petocz, P. and Alfano, M. (2024), "How should we change teaching and assessment in response to increasingly powerful generative artificial intelligence? Outcomes of the ChatGPT teacher survey", *Education and Information Technologies*, Vol. 29 No. 12, pp. 15403-15439.
- Chan, C.K.Y. (2023), "A comprehensive AI policy education framework for university teaching and learning", *International Journal of Educational Technology in Higher Education*, Vol. 20 No. 1, p. 38.
- Chin, W.W. (1998), "Commentary: issues and opinion on structural equation modeling", *MIS Quarterly*, Vol. 22 No. 1, pp. vii-xvi.
- Cohen, J. (1998), "Set correlation and contingency tables", *Applied Psychological Measurement*, Vol. 12 No. 4, pp. 425-434.
- Choi, J.I., Yang, E. and Goo, E.H. (2024), "The effects of an ethics education program on artificial intelligence among middle school students: analysis of perception and attitude changes", *Applied Sciences*, Vol. 14 No. 4, p. 1588.

-
- Collie, R.J. and Martin, A.J. (2024), "Teachers' motivation and engagement to harness generative AI for teaching and learning: the role of contextual, occupational, and background factors", *Computers and Education: Artificial Intelligence*, Vol. 6, p. 100224.
- Corbin, T., Bearman, M., Boud, D. and Dawson, P. (2025), "The wicked problem of AI and assessment", *Assessment and Evaluation in Higher Education*, pp. 1-17.
- Dabis, A. and Csáki, C. (2024), "AI and ethics: investigating the first policy responses of higher education institutions to the challenge of generative AI", *Humanities and Social Sciences Communications*, Vol. 11 No. 1, pp. 1-13.
- Dawson, P. (2020), *Defending Assessment Security in a Digital World: Preventing e-Cheating and Supporting Academic Integrity in Higher Education*, Routledge.
- Deci, E.L. and Ryan, R.M. (1985), "The general causality orientations scale: self-determination in personality", *Journal of Research in Personality*, Vol. 19 No. 2, pp. 109-134.
- Đerić, E., Frank, D. and Vuković, D. (2025), "Exploring the ethical implications of using generative AI tools in higher education", *Informatics*, MDPI, Vol. 12 No. 2, p. 36.
- Eaton, S.E., Moya Figueroa, B.A., McDermott, B., Kumar, R., Brennan, R. and Wiens, J. (2025), "What should we be assessing exactly? Higher education staff narratives on gen AI integration of assessment in a postplagiarism era", *Assessment and Evaluation in Higher Education*, Vol. 51 No. 4, pp. 1-20.
- Eshet, Y. (2004), "Digital literacy: a conceptual framework for survival skills in the digital era", *Journal of Educational Multimedia and Hypermedia*, Vol. 13 No. 1, pp. 93-106.
- Foltynek, T., Bjelobaba, S., Glendinning, I., Khan, Z.R., Santos, R., Pavletic, P. and Kravjar, J. (2023), "ENAI recommendations on the ethical use of artificial intelligence in education", *International Journal for Educational Integrity*, Vol. 19 No. 1, pp. 1-4.
- Foung, D., Lin, L. and Chen, J. (2024), "Reinventing assessments with ChatGPT and other online tools: opportunities for GenAI-empowered assessment practices", *Computers and Education: Artificial Intelligence*, Vol. 6, p. 100250.
- Furze, L., Perkins, M., Roe, J. and MacVaugh, J. (2024), "The AI assessment scale (AIAS) in action: a pilot implementation of GenAI-supported assessment", *Australasian Journal of Educational Technology*, Vol. 40 No. 4, pp. 38-55.
- García-López, I.M. and Trujillo-Liñán, L. (2025), "Generative artificial intelligence in education: ethical challenges, regulatory frameworks and educational quality in a systematic review of the literature", *Frontiers in Education*, Vol. 10, p. 1565938.
- Giannakos, M., Azevedo, R., Brusilovsky, P., Cukurova, M., Dimitriadis, Y., Hernandez-Leo, D., ... Rienties, B. (2024), "The promise and challenges of generative AI in education", *Behaviour and Information Technology*, pp. 1-27.
- Gilster, P. (1997), *Digital Literacy*, John Wiley and Sons.
- Hair, J. and Alamer, A. (2022), "Partial least squares structural equation modeling (PLS-SEM) in second language and education research: guidelines using an applied example", *Research Methods in Applied Linguistics*, Vol. 1 No. 3, p. 100027.
- Hair, J., Hollingsworth, C.L., Randolph, A.B. and Chong, A.Y.L. (2017), "An updated and expanded assessment of PLS-SEM in information systems research", *Industrial Management and Data Systems*, Vol. 117 No. 3, pp. 442-458.
- Hair, J.F., Risher, J.J., Sarstedt, M. and Ringle, C.M. (2019), "When to use and how to report the results of PLS-SEM", *European Business Review*, Vol. 31 No. 1, pp. 2-24.
- Hair, J.F., Hult, G.T.M., Ringle, C.M., Sarstedt, M., Danks, N.P. and Ray, S. (2021), "Partial least squares structural equation modeling (PLS-SEM) using R: a workbook", Springer International Publishing.
- He, G. (2025), "Predicting learner autonomy through AI-supported self-regulated learning: a social cognitive theory approach", *Learning and Motivation*, Vol. 92, p. 102195.

- He, L. and Li, C. (2023), "Continuance intention to use mobile learning for second language acquisition based on the technology acceptance model and self-determination theory", *Frontiers in Psychology*, Vol. 14, p. 1185851.
- Henseler, J., Ringle, C.M. and Sarstedt, M. (2015), "A new criterion for assessing discriminant validity in variance-based structural equation modeling", *Journal of the Academy of Marketing Science*, Vol. 43 No. 1, pp. 115-135.
- Holmes, W., Porayska-Pomsta, K., Holstein, K., Sutherland, E., Baker, T., Shum, S.B., ... Koedinger, K. R. (2022), "Ethics of AI in education: towards a community-wide framework", *International Journal of Artificial Intelligence in Education*, Vol. 32 No. 3, pp. 1-23.
- Hu, L.T. and Bentler, P.M. (1999), "Cutoff criteria for fit indexes in covariance structure analysis: Conventional criteria versus new alternatives", *Structural Equation Modeling: A Multidisciplinary Journal*, Vol. 6 No. 1, pp. 1-55.
- Huang, L. (2023), "Ethics of artificial intelligence in education: student privacy and data protection", *Science Insights Education Frontiers*, Vol. 16 No. 2, pp. 2577-2587.
- Huang, K., Huang, J. and Catteddu, D. (2024), "GenAI data security", *Generative AI Security: Theories and Practices*, Springer Nature Switzerland, Cham, pp. 133-162.
- Izevbogie, B., Oboh, G. and Okonkwo, R. (2025), "The ethical use of generative AI in the Nigerian higher education sector", *World Journal of Advanced Research and Reviews*, Vol. 20 No. 2, pp. 123-130, available at: <https://journalwjarr.com/content/ethical-use-generative-ai-nigerian-higher-education-sector>.
- Joseph, F., Tomas, G., Hult, M. and Christian, M. (2017), *A Primer on Partial Least Squares Structural Equation Modeling (PLS-SEM)*, Library of Congress Cataloging-in-Publication Data.
- Jukiewicz, M. (2024), "The future of grading programming assignments in education: the role of ChatGPT in automating the assessment and feedback process", *Thinking Skills and Creativity*, Vol. 52, p. 101522.
- Khlaif, Z.N., Alkouk, W.A., Salama, N. and Abu Eideh, B. (2025), "Redesigning assessments for AI-enhanced learning: a framework for educators in the generative AI era", *Education Sciences*, Vol. 15 No. 2, p. 174.
- Kim, J., Yu, S., Detrick, R. and Li, N. (2025), "Exploring students' perspectives on generative AI-assisted academic writing", *Education and Information Technologies*, Vol. 30 No. 1, pp. 1265-1300.
- Law, L. (2024), "Application of generative artificial intelligence (GenAI) in language teaching and learning: a scoping literature review", *Computers and Education Open*, Vol. 6 No. 2, p. 100174.
- Lazăr, A.M., Repanovici, A., Popa, D., Ionas, D.G. and Dobrescu, A.I. (2024), "Ethical principles in AI use for assessment: exploring students' perspectives on ethical principles in academic publishing", *Education Sciences*, Vol. 14 No. 11, p. 1239.
- Lim, T., Gottipati, S. and Cheong, M. (2025), "What students really think: unpacking AI ethics in educational assessments through a triadic framework", *International Journal of Educational Technology in Higher Education*, Vol. 22 No. 1, p. 56.
- Luo, J. (2024), "A critical review of GenAI policies in higher education assessment: a call to reconsider the "originality" of students' work", *Assessment and Evaluation in Higher Education*, Vol. 49 No. 5, pp. 651-664.
- Luo, J. (2025), "How does GenAI affect trust in teacher-student relationships? Insights from students' assessment experiences", *Teaching in Higher Education*, Vol. 30 No. 4, pp. 991-1006.
- Magrill, J. and Magrill, B. (2024), "Preparing educators and students at higher education institutions for an AI-driven world", *Teaching and Learning Inquiry*, Vol. 12, pp. 1-9.
- Malik, A.R., Pratiwi, Y., Andajani, K., Numertayasa, I.W., Suharti, S. and Darwis, A. (2023), "Exploring artificial intelligence in academic essay: higher education student's perspective", *International Journal of Educational Research Open*, Vol. 5, p. 100296.

-
- Marín, Y.R., Caro, O.C., Rituay, A.M.C., Llanos, K.A.G., Perez, D.T., Bardales, E.S., ... Santos, R.C. (2025), "Ethical challenges associated with the use of artificial intelligence in university education", *Journal of Academic Ethics*, Vol. 23 No. 4, pp. 2443-2467.
- Memarian, B. and Doleck, T. (2023), "Fairness, accountability, transparency, and ethics (FATE) in artificial intelligence (AI) and higher education: a systematic review", *Computers and Education: Artificial Intelligence*, Vol. 5, p. 100152.
- Moorhouse, B.L., Yeo, M.A. and Wan, Y. (2023), "Generative AI tools and assessment: guidelines of the world's top-ranking universities", *Computers and Education Open*, Vol. 5, p. 100151.
- Ng, W. (2012), "Can we teach digital natives digital literacy?", *Computers and Education*, Vol. 59 No. 3, pp. 1065-1078.
- Nguyen, K.V. (2025), "The use of generative AI tools in higher education: ethical and pedagogical principles", *Journal of Academic Ethics*, Vol. 23 No. 3, pp. 1-21.
- Oved, O. and Alt, D. (2025), "Teachers' technological pedagogical content knowledge (TPACK) as a precursor to their perceived adopting of educational AI tools for teaching purposes", *Education and Information Technologies*, Vol. 30 No. 10, pp. 1-27.
- Rabotapi, T. and Matope, S. (2024), "A dataset on WhatsApp groups effectiveness in inducting first years to university", *Data in Brief*, Vol. 54, p. 110456.
- Rasul, T., Nair, S., Kalendra, D., Balaji, M.S., de Oliveira Santini, F., Ladeira, W.J., ... Hossain, M.U. (2024), "Enhancing academic integrity among students in GenAI era: a holistic framework", *The International Journal of Management Education*, Vol. 22 No. 3, p. 101041.
- Roe, J., Perkins, M. and Ruelle, D. (2024), "Understanding student and academic staff perceptions of AI use in assessment and feedback", arXiv preprint arXiv:2406.15808.
- Ryan, R.M. and Deci, E.L. (2000), "Intrinsic and extrinsic motivations: classic definitions and new directions", *Contemporary Educational Psychology*, Vol. 25 No. 1, pp. 54-67.
- Ryan, R.M. and Moller, A.C. (2017), "Competence as central, but not sufficient, for high-quality motivation", *Handbook of Competence and Motivation: Theory and Application*, Vol. 2, pp. 216-238.
- Shailendra, S., Kadel, R. and Sharma, A. (2024), "Framework for adoption of generative artificial intelligence (GenAI) in education", *IEEE Transactions on Education*, Vol. 67 No. 5.
- Shin, D. (2021), "The effects of explainability and causability on perception, trust, and acceptance: Implications for explainable AI", *International Journal of Human-Computer Studies*, Vol. 146, p. 102551.
- Shuford, J. (2024), "Examining ethical aspects of AI: addressing bias and equity in the discipline", *Journal of Artificial Intelligence General Science (JAIGS)*, ISSN: 3006-4023, Vol. 3 No. 1, pp. 262-280.
- Szabó, F. and Szoke, J. (2024), "How does generative AI promote autonomy and inclusivity in language teaching?", *ELT Journal*, Vol. 78 No. 4, pp. 478-488.
- Thulasiram, P.P. (2025), "Explainable artificial intelligence (XAI): enhancing transparency and trust in machine learning model", *International Journal for Innovative Engineering and Management Research*, Vol. 14 No. 1, pp. 204-213.
- Turner, P. (2025), "Mapping a multidimensional framework for GenAI in education", *EDUCAUSE Review* (Online).
- Wiese, L.J., Patil, I., Schiff, D.S. and Magana, A.J. (2025), "AI ethics education: a systematic literature review", *Computers and Education: Artificial Intelligence*, Vol. 8, p. 100405.
- Xia, Q., Weng, X., Ouyang, F., Lin, T.J. and Chiu, T.K. (2024), "A scoping review on how generative artificial intelligence transforms assessment in higher education", *International Journal of Educational Technology in Higher Education*, Vol. 21 No. 1, p. 40.
- Yim, I.H.Y. and Su, J. (2025), "Artificial intelligence (AI) learning tools in K-12 education: a scoping review", *Journal of Computers in Education*, Vol. 12 No. 1, pp. 93-131.
- Zhang, C. and Magerko, B. (2025), "Generative AI literacy: a comprehensive framework for literacy and responsible use", arXiv preprint arXiv:2504.19038.

- Zhuang, M., Long, S., Martin, F. and Castellanos-Reyes, D. (2025), "The affordances of artificial intelligence (AI) and ethical considerations across the instruction cycle: a systematic review of AI in online higher education", *The Internet and Higher Education*, Vol. 67, p. 101039.
- Zlotnikova, I., Hlomani, H., Mokgetse, T. and Bagai, K. (2025), "Establishing ethical standards for GenAI in university education: a roadmap for academic integrity and fairness", *Journal of Information, Communication and Ethics in Society*, Vol. 23 No. 2, pp. 188-216.

Further reading

- Cao, Y., Li, S., Liu, Y., Yan, Z., Dai, Y., Yu, P.S. and Sun, L. (2023), "A comprehensive survey of ai-generated content (AIGC): a history of generative AI from Gan to ChatGPT", arXiv preprint arXiv:2303.04226.
- Ning, Y., Teixayavong, S., Shang, Y., Savulescu, J., Nagaraj, V., Miao, D., Wang, J., Cao, Q. and Liu, N. (2024), "Generative artificial intelligence and ethical considerations in health care: a scoping review and ethics checklist", *The Lancet Digital Health*, Vol. 6 No. 11, pp. e848-e856.
- Perkins, M., Furze, L., Roe, J. and MacVaugh, J. (2024), "The artificial intelligence assessment scale (AIAS): a framework for ethical integration of generative AI in educational assessment", *Journal of University Teaching and Learning Practice*, Vol. 21 No. 6, pp. 49-66.

Appendix

Table A1. Questionnaire items

Constructs	Definitions	Items	References
GenAI transparency (GT)	“The degree to which students understand how GenAI tools operate and produce outputs”	- “I am aware of how GenAI tools generate their content - I understand the limitations of GenAI outputs - I believe GenAI-generated content should be clearly disclosed as AI-generated - I seek information about how GenAI systems function before using them”	(Shin, 2021)
GenAI autonomy (GA)	“The extent to which students independently verify and make decisions about GenAI outputs without over-relying on AI”	- “I verify information generated by GenAI before using it in my work - I make independent judgments about content generated by GenAI - I use GenAI as a support tool in my work - I recognize when human input is needed to improve GenAI suggestions”	(Collie and Martin, 2024)
GenAI privacy (GP)	“The concern for protecting personal and academic data when using GenAI tools”	- “I am cautious about sharing sensitive information with GenAI tools - I understand the privacy risks of using GenAI in academic settings - I prefer GenAI tools that protect data confidentiality - I review privacy policies before using new GenAI platforms”	(Baig and Yadegaridehkordi, 2025)
GenAI ethics literacy (GEL)	“The level of knowledge and understanding students have about ethical issues in GenAI use”	- “I am familiar with ethical concerns related to using GenAI in education - I have been exposed to guidelines about the ethical use of GenAI tools - I can identify unethical uses of GenAI in academic assessments - I feel confident making ethical decisions when using GenAI”	(Huang, 2023)
GenAI policy clarity (GPC)	“The degree to which students perceive university policies on GenAI use as clear and understandable”	- “My university has clear policies about the acceptable use of GenAI - I understand the rules regarding GenAI use in assessments - My university clearly communicates the consequences of unethical GenAI use. d - I receive sufficient guidance on when and how GenAI tools can be used”	(Chan, 2023)
GenAI fairness (GF)	“The extent to which students perceive that GenAI use in assessments is treated fairly and equitably”	- “I believe GenAI policies ensure fairness among students - I believe using GenAI tools does not give certain students an unfair advantage - I believe GenAI usage rules are applied consistently across courses - I believe assessment practices address ethical concerns related to GenAI content”	(Shin, 2021)

(continued)

Table A1. Continued

Constructs	Definitions	Items	References
Ethical GenAI use in assessments (EGA)	“The degree to which students engage in responsible and ethical behavior when using GenAI in academic tasks”	“-I use GenAI tools in ways that align with academic integrity principles -I avoid plagiarism when using GenAI outputs -I disclose the use of GenAI to assist in my academic assignments -I check that GenAI-generated content meets ethical standards before submitting it”	(Baig and Yadegaridehkordi, 2025)

Corresponding author

Elaheh Yadegaridehkordi can be contacted at: yellahe@gmail.com