



Behavioral and ethical predictors of continuous intention to use generative AI responsibly in higher education

Ibrahim Arpaci ^{1,2}

 0000-0001-6513-4569

Mustafa Baloglu ^{3*}

 0000-0003-1874-9004

¹ Department of Computer Engineering, Faculty of Engineering, Bursa Uludag University, Bursa, TURKEY

² Department of Computer Science and Engineering, College of Informatics, Korea University, Seoul, SOUTH KOREA

³ Department of Special and Gifted Education, College of Education, United Arab Emirates University, Al Ain, UNITED ARAB EMIRATES

* Corresponding author: baloglu@hotmail.com

Citation: Arpaci, I., & Baloglu, M. (2026). Behavioral and ethical predictors of continuous intention to use generative AI responsibly in higher education. *Contemporary Educational Technology*, 18(3), Article ep673. <https://doi.org/10.30935/cedtech/18951>

ARTICLE INFO

Received: 7 May 2026

Accepted: 6 Jul 2026

ABSTRACT

This study examined factors that support the sustainable and responsible integration of generative artificial intelligence (GenAI) into higher education. The research model was based on the theory of planned behavior (TPB) and was extended to include ethical considerations, such as authenticity, originality, responsibility, and confidentiality. In the qualitative phase, structured interviews with faculty members across various departments examined the potential benefits, limitations, and ethical concerns of GenAI use in higher education. The qualitative findings informed the development of the theoretical framework and the research model. In the quantitative phase, PLS-SEM was used to test the model with data from 1,261 GenAI users. The results showed that authenticity, originality, responsibility, and confidentiality significantly predicted attitudes (ATs) toward GenAI, while ATs, perceived behavioral control, and subjective norms significantly predicted continuous intention. The findings contribute by testing an ethically grounded extension of the TPB for the responsible integration of GenAI in higher education. They also emphasize the need for clear behavioral rules and ethical guidelines, developed with relevant stakeholders, to support sustainable and responsible use of GenAI.

Keywords: artificial intelligence, AI ethical use, sustainability, mixed methods

INTRODUCTION

Generative artificial intelligence (GenAI) tools are advanced technologies powered by artificial intelligence (AI), engineered for interactive dialogue with users using natural language (McTear, 2021). AI tools utilize machine learning algorithms and large language processing techniques to interpret inputs and generate content based on extensive existing data. Therefore, these tools can produce biased content, reflecting biases in the underlying algorithms or the training data (Arpaci, 2024).

GenAI has become increasingly valuable in academia and higher education because it offers rapid access to information and helps reduce the time and effort required for academic tasks (Belkina et al., 2025; Kasneci et al., 2023). By entering a simple prompt, researchers and students can obtain timely and largely relevant responses that support a wide range of learning and research activities (Gibreel & Arpaci, 2025). This timely assistance and ubiquitous access to information accelerate research and help researchers stay up-to-date with the latest advancements in their respective fields (Kumar et al., 2024; Nguyen et al., 2022).

To ensure the sustainable and responsible use of AI in higher education and academic research settings, ethical issues must be addressed comprehensively (Dave et al., 2023; Mhlanga, 2023). When using GenAI in

research, it is essential to consider the behavioral and ethical issues that may arise (King, 2023; Zhou et al., 2023). This study primarily focuses on the responsible integration of GenAI in academia and higher education. Therefore, the research question addressed in this study is: "How do behavioral and ethical factors influence the continuous intention to integrate GenAI in higher education responsibly?" Responsible integration refers to the sustainable and ethical implementation of GenAI in academia and higher education (Arpaci et al., 2025). Sustainable implementation aims to ensure that these systems are used efficiently while addressing ethical issues, long-term impacts, and the overall well-being of the academic community.

LITERATURE REVIEW

GenAI technologies have the power to revolutionize learning and teaching techniques by enabling effective knowledge management (Storey, 2024). Researchers can use these systems to find relevant resources, key research articles, and literature recommendations (Kim et al., 2024). However, GenAI can hallucinate and produce articles that do not exist, and therefore, it is not recommended to use such systems to conduct a literature review (Hwang & Jeong, 2025). GenAI tools are currently in development and are expected to make significant strides in the years to come.

GenAI can support scholarly publishing by providing academic support services (Ganguly et al., 2025). These systems can assist in performing routine tasks, such as proof-reading, summarizing, paraphrasing, checking language, grammar, and spelling in articles, as well as translating texts (Söllner et al., 2025). This minimizes the burden on academic staff and provides them with accurate and timely assistance (Stokel-Walker & Van Noorden, 2023).

Researchers and students can use these systems to develop new approaches and gain diverse perspectives on their subjects of study. GenAI can provide immediate access to a vast repository of information across diverse domains, enabling researchers to explore complex topics by generating explanations, summaries, and insights from multiple perspectives (Lund & Wang, 2023). GenAI can also serve as a research assistant, generating ideas, diverse perspectives, and brainstorming insights (Bouschery et al., 2023; English et al., 2025).

AI systems provide a supportive environment for language development by offering immediate feedback and constructive corrections (Kuhail et al., 2023). By interacting with these systems, language learners can improve their grammar, vocabulary, and writing skills (Kohnke et al., 2023). This interactivity improves language proficiency and increases communication confidence (Huang et al., 2022; Yeh, 2025).

Another unique benefit of GenAI is personalized, lifelong learning (Chiu et al., 2023; Vorobyeva et al., 2025). These tools can adapt to learning preferences and styles (Shen et al., 2023). Educators and students can access personalized recommendations, individualized study materials, and adaptive learning environments (Farrokhnia et al., 2023). Thus, they can have an interactive and real-time learning experience that promotes learning performance, critical thinking, higher-order thinking, problem-solving skills, and active student engagement (Chen et al., 2023; Wang & Fan, 2025).

On the other hand, it is essential to know that GenAI can have potential disadvantages and limitations (Zhou et al., 2024). A significant drawback of these tools is their lack of common sense or contextual awareness, which can result in incomprehensible or irrelevant responses (Hutson, 2021). Notably, these systems lack the emotional sensitivity and insight of humans (Hill et al., 2015). Furthermore, due to biases in training data or the algorithms employed, these systems may produce inaccurate or biased answers (Bender et al., 2021). For example, they may present a combination of real and fabricated articles or references (Alkaissi & McFarlane, 2023). This raises concerns that the use of these technologies in scientific writing could compromise academic honesty, credibility, and accuracy (Qin et al., 2020).

Consequently, the use of AI in higher education and academia raises numerous ethical concerns, including concerns about scientific integrity (Cotton et al., 2023; Gašević et al., 2023). In addition, ethical issues related to information security, dependency, privacy, and equal access are also essential factors to consider (Nguyen, 2025). Recognizing these challenges is crucial for utilizing GenAI responsibly and effectively in academia and higher education.

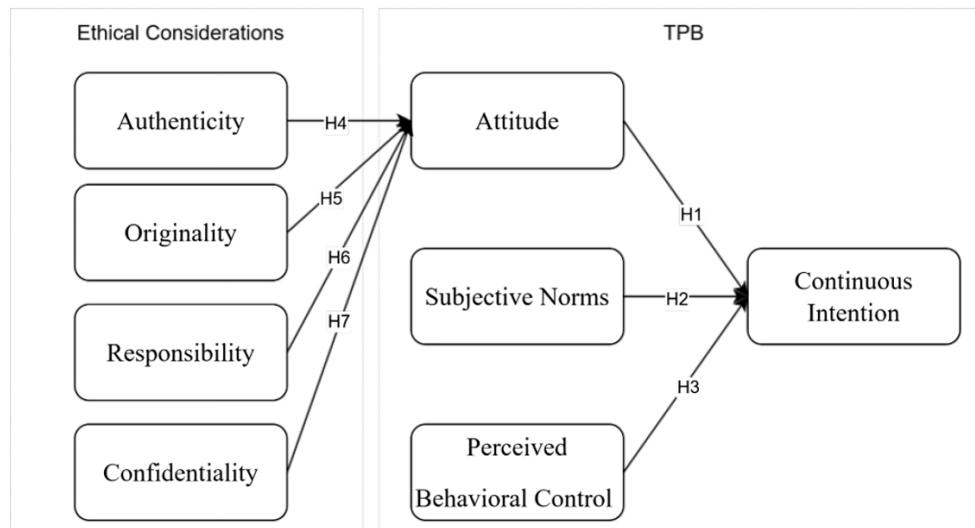


Figure 1. Research model (developed by the authors)

While existing studies have examined AI adoption in higher education through behavioral frameworks, the specific intersection of ethical constructs and sustainable integration remains underexplored. The present study addresses this gap by developing and testing an ethically grounded predictive model tailored to the higher education context, thereby extending the theoretical scope of prior work and contributing a more comprehensive account of the factors shaping responsible AI integration in academia.

The contribution of the present model lies not in applying the theory of planned behavior (TPB) alone, but in specifying how ethical concerns identified in the qualitative phase operate as antecedents of attitude (AT) within a continuous-use framework. This positioning distinguishes the study from TPB-based AI adoption research that treats behavioral intention primarily in terms of ATs, norms, and perceived behavioral control (PBC), without modeling ethical concerns as construct-level predictors.

RESEARCH MODEL AND HYPOTHESES

The TPB seeks to explain human behavior through PBC, ATs, and subjective norms (SN) (Ajzen, 2012). This theory suggests that all these constructs contribute to defining an individual's behavioral intention, i.e., the readiness and willingness to engage in a particular behavior (Ajzen, 2002). The TPB has been effectively used in previous studies that focus on identifying factors explaining the use of emerging technologies in higher education (Al-Emran et al., 2022).

This study aimed to identify the key factors predicting continuous intention to use GenAI responsibly. Although well-established theories such as UTAUT and TAM were considered, the TPB was selected as the theoretical foundation given its explicit emphasis on behavioral intention. The TPB was then expanded to incorporate four ethical constructs: authenticity, originality, responsibility, and confidentiality.

For this study, AI refers broadly to computational systems capable of performing tasks that typically require human intelligence. At the same time, GenAI refers to a specific subset of AI systems that generate novel content using large-scale language models. Confidentiality, authenticity, originality, and responsibility are treated as distinct ethical constructs: confidentiality concerns the protection of personal data and intellectual property; authenticity pertains to the accuracy and trustworthiness of AI-generated information; originality addresses the avoidance of derivative or plagiarized content; and responsibility encompasses the obligation of both users and developers to adhere to ethical standards and societal norms. [Figure 1](#) illustrates the research model.

Attitude

AT refers to a person's tendency to perceive a particular behavior either favorably or unfavorably, based on their ideas about its consequences or outcomes (Ajzen, 2012). In this study, the AT toward using GenAI is defined as a favorable outlook, rooted in the belief that utilizing GenAI is a good choice. In contrast,

“continuous intention” is defined as an enduring, determined inclination in an individual to use GenAI consistently. Previous research has affirmed that a favorable AT has a noteworthy positive impact on intentions to use GenAI for informal learning (Wu & Dong, 2025). A positive AT toward GenAI may also play a significant role in predicting individuals’ continuous intention to use these systems. Thus,

H1. ATs significantly predict continuous intention to use GenAI.

Subjective Norms

SN refer to “the perceived societal pressure to engage or refrain from engaging in a specific behavior.” (Ajzen, 1991, p. 195). Previous findings have consistently shown the significant and positive impact of SN on intentions to use various technologies, including WhatsApp stickers (Al-Marroof et al., 2021), smartwatches (Al-Emran et al., 2023), mobile learning (Al-Emran et al., 2020), mobile cloud services (Arpaci, 2019), and GenAI (Al-Qaysi et al., 2025). Therefore,

H2. SN are a significant predictor of continuous intention.

Perceived Behavioral Control

PBC is defined as “an individual’s perception of the ease or difficulty associated with performing a specific behavior” (Ajzen, 1991, p. 183). The PBC is a significant determinant of intentions to use innovative technologies, such as smartwatches (Al-Emran et al., 2023), and wearable technologies (Al-Emran et al., 2022). PBC includes individuals’ beliefs and self-assessments of their abilities, knowledge, control, and confidence in decision-making about GenAI use. Therefore, it is likely to be associated with continued GenAI use,

H3. PBC is a significant predictor of continuous intention.

Authenticity

Authenticity is defined through goodwill, practical wisdom, virtuousness, and an Aristotelian conception of ethos that includes Fromm’s concepts of the false self and the true self (Deptula et al., 2025). This quality inspires confidence in the accuracy and authenticity of the information produced by ensuring consistency, reliability, and trustworthiness (Arpaci & Sevinc, 2022). The acceptability of the content created by AI tools depends on their accuracy, reliability, and authenticity (Maras & Alexandrou, 2019). A study evaluating the authenticity of responses from ChatGPT-3.5 and ChatGPT-4 found that the recent version produced more authentic responses than the previous one (Elkhataf, 2023). As AI systems evolve, the accuracy, reliability, and authenticity of the information they generate increase, which positively influences ATs towards GenAI use. Therefore,

H4. Authenticity is a significant predictor of ATs toward GenAI use.

Originality

In the context of AI, originality refers to the avoidance of unintentional copying of copyrighted content (Gaffar & Albarashdi, 2024). Recent findings highlight AI tools’ remarkable ability to produce complex text that is undetectable by plagiarism-detection software, indicating their high potential to generate highly original content similar to that of a human author (Khalil & Er, 2023). However, overreliance on such systems without verification can lead to significant errors or a lack of originality in the content generated (Carvalho & Ivanov, 2024). Confirming this, previous findings have shown that AI tools tend to reproduce sections of articles from online sources without proper references (AlAfnan et al., 2023). The extent to which GenAI generates distinct, non-derivative content significantly influences users’ ATs (Hwang et al., 2025). When AI-generated content demonstrates originality by not overtly copying or imitating existing work, users’ AT toward GenAI use can be positively influenced (Ali & Arpaci, 2025). Therefore,

H5. Originality is a significant predictor of ATs toward GenAI use.

Responsibility

GenAI is viewed as a socio-technological system where human actors, AI technologies, institutional policies, data infrastructures, and operational processes all evolve together in an interconnected manner (Janssen, 2025). Therefore, there is a need for shared responsibility and accountability, where the parties

involved work together (Ulnicane, 2025). Responsibility, which entails an obligation, duty, or accountability to adhere to ethical principles, societal norms, and the welfare of individuals and communities, is a pivotal subject of discussion within the domain of artificial agents (Johnson & Noorman, 2014). The lack of clarity in assigning responsibility can give rise to legal and ethical concerns (Atkins et al., 2021). Thus, establishing clear boundaries and rules is crucial for appropriately assigning liability and protecting users when employing GenAI (Wang et al., 2023). Users' development of positive ATs towards GenAI use depends on their belief that AI systems and service providers act responsibly and ethically, consider potential adverse effects such as bias, and comply with ethical standards. Therefore,

H6. Responsibility is a significant predictor of ATs toward the use of GenAI.

Confidentiality

Confidentiality in the context of GenAI is fundamentally about protecting confidential information, interactions, privacy, and intellectual property (Prinz, 2025). Previous research has emphasized that perceived security, encompassing confidentiality, integrity, and availability, significantly predicts ATs toward using cloud services (Arpaci, 2019). Confidentiality issues may arise when processing personal data on AI-driven platforms; therefore, it is critical to develop robust data protection protocols to ensure adequate privacy protection (Adhikari et al., 2024; Sharma & Sharma, 2023). These platforms may also face accusations of compromising data privacy and generating false information (Chowdhury et al., 2023). Users' trust in GenAI tools to protect their privacy and confidential information would positively influence their ATs and willingness to participate in the system. Therefore,

H7. Confidentiality is a significant predictor of ATs toward the use of GenAI.

METHODOLOGY

Research Design

This study employed an exploratory sequential mixed-methods design with an instrument development component. In the first phase, qualitative data were collected and analyzed to identify the perceived benefits, limitations, and ethical concerns associated with the use of GenAI in higher education. These findings informed both the conceptual specification of the research model and the development of items representing the ethical constructs for the quantitative phase. In the second phase, the themes that emerged from the qualitative inquiry were translated into measurable indicators and quantitatively examined using exploratory factor analysis (EFA) and PLS-SEM. Thus, the mixed-methods design was structured sequentially, with the qualitative phase directly informing scale development and subsequent model testing.

In the qualitative phase, faculty members from different academic disciplines were interviewed using a structured set of research questions. The interviews were analyzed to identify recurring themes in participants' views on integrating GenAI into academia and higher education. The interview questions addressed three main issues: the benefits of integrating GenAI into academia and higher education, the primary limitations of its use in these settings, and the ethical concerns and challenges arising from its integration into research and academic practices.

Sample and Procedure

The initial phase of the study included a sample of twenty faculty members from governmental universities located in Turkey. Among the interviewees, 13 were male, and 7 were female ($n = 20$; mean age = 32.05). Furthermore, the faculty members were categorized as follows: 11 research assistants, 5 assistant professors, 2 lecturers, 1 associate professor, and 1 full professor. The recruitment of participants encompassed four disciplines representing related fields of study: engineering (e.g., computer, software, industry, and architecture) ($n = 13$), social sciences (e.g., psychology and politics) ($n = 3$), educational sciences (e.g., foreign languages and literacy) ($n = 2$), and natural sciences (e.g., mathematics and physics) ($n = 2$). The study employed a purposive sampling strategy, targeting faculty members with diverse perspectives, expertise, and experience with GenAI. Among the participants, 35% reported using it occasionally (once a week), 20% stated

using it very frequently (every day), 20% mentioned never using it, 15% indicated frequent usage (4-5 times a week), and 10% reported rare usage (once a month).

In the second phase of the study, 1,629 individuals participated with a mean age of 23.25 (standard deviation = 7.397; age range = 18-50). The majority of them were university students, including first-year students (39.3%), sophomores (25.9%), juniors (19.8%), seniors (11.5%), MSc. students (3.2%), and PhD students (0.3%). They represented diverse departments, including engineering, medicine, law, and education. Among the participants, 51.07% were female (832 females) and 48.93% were male (797 males). Additionally, 77.4% of participants reported using a GenAI tool, such as ChatGPT, ChatSonic, Copilot, YouChat, or Jasper. Therefore, 368 non-users were excluded from the primary analysis, leaving 1,261 actual users for further study. Accordingly, the quantitative phase focused specifically on participants with prior experience using GenAI, as the model examined continuous intention rather than initial adoption.

A sample of 200 to 500 respondents is generally considered sufficient for SEM (Tabachnick & Fidell, 2012); the present sample of 1,261 users exceeds this range, so the sample size was deemed appropriate. The study recruited volunteers who expressed interest in participating. Before their involvement, individuals were informed about the research objectives and received comprehensive information, ensuring that the data collected would be used exclusively for research purposes. The study strictly adhered to ethical guidelines to protect participant confidentiality and data security and received approval (#2023-7) before data collection. Participants provided informed consent and retained the right to withdraw at any time without consequence.

Data Collection

The online data collection tool used in the study consisted of a questionnaire comprising eight scales and questions on demographic information and usage behavior. A Likert scale, ranging from "1 = strongly disagree" to "5 = strongly agree," was utilized to assess eight factors with 26 items (see [Appendix A](#)). The items measuring AT (Alalwan et al., 2018), SN (Al-Emran et al., 2020), PBC (Ajzen, 2002), and continuous intention (Al-Sharafi et al., 2023) were adapted from previous studies. The items measuring authenticity, originality, responsibility, and confidentiality were developed based on the qualitative phase findings and subsequently evaluated in the current study.

Content Validity

The study used a structured coding process to ensure rigor and validity, including thematic analysis, double coding, iterative refinement, expert validation, inter-rater reliability assessment, member checking, and theoretical fit. To ensure reliability and minimize bias, a double-coding approach was used, with two independent researchers coding the transcripts independently. In addition, Cohen's kappa ($\kappa = 0.80$) was calculated to assess inter-rater reliability. The items representing the ethical constructs were derived directly from recurrent qualitative codes and themes, and the resulting item pool was reviewed for clarity and construct relevance, refined before pilot administration, and retained or revised during subsequent validation. Thus, qualitative themes guided both the development of the research model and the construction of the ethics-related measures.

Quantitative Data Analysis

Because some of the study constructs were newly developed, the quantitative analysis proceeded in two stages. First, SPSS (Version 29) was used to conduct EFA to examine the factor structure of the newly developed items. Furthermore, EFA was applied to the full scale to examine the instrument's structure. Therefore, the reported measurement results include both the newly developed ethical constructs and the adapted TPB constructs. Second, PLS-SEM was performed in SmartPLS (version 4) to assess the measurement model and test the structural relationships among the constructs. This sequence allowed the newly developed measures to be examined before the hypothesized paths were evaluated.

Table 1. Themes and categories

Theme	Category	Sample statements
Advantages	Speed and efficiency (f = 10/20)	"It can provide faster literature search capabilities," "It reduces the time I spend on tasks," "It allows quick access to a much wider range of resources," "It accelerates research," "and provides fast and accurate information compared to searching on Google."
	Assistance and support (f = 4/20)	"Facilitating routine tasks, providing translation and correction services," "Being highly effective in editing scientific texts," and "Being able to provide immediate answers when we get stuck."
	Versatility and creativity (f = 11/20)	"Being able to produce different outputs usable in multiple fields," "Providing new perspectives and enhancing imagination," "Being able to provide detailed answers and further elaborate on them," and "Being knowledgeable in almost every field."
	Faster access to information resources (f = 6/20)	"Enabling faster access to resources," "Being able to access information on unfamiliar topics within seconds instead of searching the Internet," "It speeds up research and allows for time-saving by gaining an overview of the subject."
	Different perspectives and insights (f = 6/20)	"Providing a more comprehensive perspective," "Very useful for figuring out the overall outline of my research," "It can provide different perspectives," and "I find it useful for contributing to the development of new methods."
Technical limitations	Facilitating routine tasks and language support (f = 13/20)	"It can also be beneficial in repetitive routine tasks," "Especially in translating texts," "Sometimes it is useful in rewriting texts with potential plagiarism," "Checking my mistakes in the text, correcting inconsistencies if there are any in the written text, and summarizing an academic text," "Being able to complete some simple tasks, e.g., summarizing, creating bibliographies, etc."
	Cognitive processing (f = 15/20)	"Sometimes wrongly interpreting questions and refusing to answer them as unethical," "Sometimes insisting on the same mistake," "Sometimes losing context," "Sometimes it gives repetitive answers."
	Training and capabilities (f = 11/20)	"Being trained until a specific date," "Inability to generate visuals and drawings, resulting in incomplete explanations," "Sometimes getting stuck," "Not providing comments on current social issues."
Ethical concerns	Authenticity (f = 12/20)	"Providing incorrect references," "It can generate incorrect information, for example, it can miswrite a code," "Not citing sources for the information it provides or generating non-existent sources," "I am a bit hesitant about the information it provides; I feel the need for verification," "Generates wrong references and non-existent sources," "Inability to provide correct solutions to certain mathematical or logical problems," "Inability to verify sources," "Sometimes providing incorrect information."
	Originality (f = 6/20)	"I frequently encounter plagiarism," "I doubt whether the answers it uses belong to copyrighted sources," and "I do not know whether the information is stolen from a source."
	Responsibility (f = 11/20)	"Having articles or assignments entirely done by GenAI," "Misleading journals by reformatting a study with GenAI," "Exploiting GenAI for unethical purposes, such as seeking ideas on unethical topics," "Students can deceive us by generating texts that are not their own using this system," "Making fake and untruthful publications."
	Confidentiality (f = 6/20)	"The possibility of my extensively generated data being obtained and used in some way," "Being used in social engineering," "Concern about presenting similar answers to other researchers working on the same topic," "Losing the confidentiality of information related to original research conducted on a unique and previously unexplored topic," "Being customized and used for unethical purposes such as hacking, password cracking, word list generation, etc."

RESULTS

Qualitative Findings

Qualitative data were collected through structured interviews lasting approximately 30 to 45 minutes each, audio-recorded with participants' consent, and concluded after the 20th interview, upon reaching theoretical saturation. Thematic analysis proceeded through inductive open coding, hierarchical organization of recurring patterns into subthemes and themes, and independent verification of inter-rater agreement by a second researcher. The statements were categorized into three overarching themes: Advantages, technical limitations, and ethical concerns of GenAI. Subsequently, the themes were reviewed to verify the relationships among the responses and to confirm the integrity and meaningfulness of the thematic groupings. **Table 1** presents the themes and categories identified through thematic analysis.

Table 2. Descriptive statistics

Factors	Mean	Standard deviation	Skewness	Kurtosis
Authenticity	9.5328	2.99277	-.050	-.323
Originality	10.2959	2.69491	-.166	-.136
Responsibility	14.0988	3.68634	-.152	-.392
Confidentiality	10.6703	2.82864	-.141	-.499
AT	13.8441	3.83604	-.236	-.217
SN	9.5856	3.10972	-.092	-.425
PBC	10.0516	2.95935	-.145	-.392
Continuous intention	10.2192	3.09203	-.232	-.405

The qualitative analysis highlighted numerous benefits that GenAI tools offer researchers, including enhanced speed and efficiency, valuable assistance and support, greater versatility and creativity in problem-solving, quicker access to information resources, diverse perspectives and insights, and the facilitation of routine tasks and language support. However, alongside these advantages, certain technical limitations were identified, notably in cognitive processing, training, and capabilities. Furthermore, ethical concerns associated with GenAI, including authenticity, responsibility, confidentiality, and originality, have been identified.

Items measuring the four ethical constructs, namely authenticity, originality, responsibility, and confidentiality, were generated inductively from recurring themes identified in the qualitative phase, reviewed by a panel of subject-matter experts for content relevance and representativeness, and subsequently refined through iterative feedback before the pilot administration. Items failing to meet the minimum factor loading threshold of .50 or showing substantively low corrected item-total correlations were candidates for removal, ensuring that only psychometrically defensible items were retained in the final instrument.

Common Method Bias

The risk of “common method bias” (CMB) associated with self-reported surveys is a significant issue across all types of research (Podsakoff et al., 2012). Since participants were asked to rate various variables on the same measurement scale in a single questionnaire, CMB may occur. Harman’s “single factor” test was employed to evaluate this concern. The results demonstrated that the first factor accounted for 19.96% of the total variance, well below the 50% threshold. This result suggests that CMB did not pose a notable issue in this study.

All variance inflation factor (VIF) values in the structural model fell within acceptable bounds, indicating no multicollinearity among the predictor constructs. In the AT equation, VIF values ranged from 1.262 for authenticity to 2.544 for responsibility, while in the continuous intention equation, VIF values were 2.007 for AT, 1.613 for PBC, and 1.569 for SN. All values remain well below the commonly applied threshold of 5.0, indicating that collinearity does not threaten the stability or interpretability of the structural estimates.

Normality

Based on the normality results, the data were normally distributed. **Table 2** presents the descriptive statistics. The skewness (standard error [SE] = .061) and kurtosis (SE = .121) values for each variable fell within the range of -3 to +3, indicating that the data were approximately normally distributed (Hair et al., 2019).

Measurement Model Assessment

The data were found to be suitable for EFA, as indicated by an excellent Kaiser-Meyer-Olkin value (KMO = .956) and a significant Bartlett’s test of sphericity, $\chi^2 (325) = 21667.097$, $p < .001$. The analysis was performed using principal component analysis with Varimax rotation and Kaiser normalization; the rotated solution converged in 7 iterations, supporting the adequacy of the extracted factor structure.

To examine construct validity at the item level, EFA was conducted as an initial screening step, given that the ethical construct items were developed inductively from the current study’s qualitative findings and had not previously been validated in this context. Factor loadings should reach at least 0.50 to demonstrate construct validity (Hair et al., 2019). The EFA results presented in **Table 3** indicate that all factor loadings exceeded the recommended threshold, supporting construct validity for the retained indicators. Moreover, the VIF scores are below the threshold, indicating no multi-collinearity among the variables. Cronbach’s alpha

Table 3. Construct validity and reliability

Factors	Items	VIF	Load	Commonality	CI-TC	AID
Authenticity ($\alpha = .816$)	A1	1.679	0.831	0.691	0.631	0.785
	A2	2.046	0.882	0.778	0.714	0.700
	A3	1.821	0.851	0.725	0.661	0.755
Originality ($\alpha = .744$)	O1	1.424	0.788	0.622	0.535	0.700
	O2	1.656	0.852	0.727	0.629	0.589
	O3	1.461	0.799	0.638	0.548	0.685
Responsibility ($\alpha = .819$)	R1	1.479	0.739	0.546	0.558	0.810
	R2	1.874	0.835	0.697	0.682	0.754
	R3	1.963	0.843	0.710	0.692	0.749
	R4	1.755	0.804	0.646	0.635	0.776
Confidentiality ($\alpha = .764$)	C1	1.664	0.846	0.717	0.628	0.644
	C2	1.599	0.831	0.691	0.604	0.672
	C3	1.444	0.794	0.630	0.554	0.728
AT ($\alpha = .850$)	AT1	1.802	0.809	0.654	0.658	0.822
	AT2	2.025	0.846	0.715	0.711	0.799
	AT3	2.029	0.845	0.715	0.710	0.800
	AT4	1.887	0.821	0.674	0.675	0.815
SN ($\alpha = .815$)	SN1	1.804	0.855	0.731	0.667	0.745
	SN2	1.780	0.852	0.725	0.662	0.751
	SN3	1.816	0.857	0.734	0.670	0.742
PBC ($\alpha = .798$)	PBC1	1.624	0.826	0.683	0.615	0.754
	PBC2	1.884	0.870	0.757	0.684	0.679
	PBC3	1.682	0.836	0.699	0.628	0.739
Continuous intention ($\alpha = .826$)	CI1	1.810	0.852	0.725	0.667	0.776
	CI2	2.009	0.877	0.769	0.709	0.734
	CI3	1.844	0.856	0.733	0.674	0.769

Note. A: Authenticity; O: Originality; R: Responsibility; C: Confidentiality; CI: Continuous intention; CI-TC: Corrected item-total correlation; & AID: Alpha if item deleted

Table 4. CR, convergent validity, and HTMT ratios

Construct	CR	AVE	CI	AT	SN	PBC	A	O	R
CI	0.827	0.614							
AT	0.850	0.587	0.918						
SN	0.815	0.595	0.722	0.713					
PBC	0.800	0.571	0.722	0.738	0.545				
A	0.785	0.646	0.579	0.590	0.590	0.463			
O	0.749	0.500	0.608	0.647	0.518	0.641	0.553		
R	0.823	0.538	0.576	0.617	0.446	0.656	0.372	0.833	
C	0.766	0.522	0.548	0.602	0.446	0.637	0.349	0.771	0.921

Note. CI: Continuous intention; A: Authenticity; O: Originality; R: Responsibility; & C: Confidentiality

values ranged from 0.744 to 0.850, indicating adequate internal consistency for each variable. These results suggest that the assessment tool's items consistently measure the same underlying construct and exhibit good reliability.

Convergent and Discriminant Validity

To assess convergent validity, average variance extracted (AVE) values were calculated to determine whether each latent variable explained more than 50% of the variance in its indicators. As shown in [Table 4](#), all AVE values met the recommended threshold of .50 (Hair et al., 2019). Composite reliability values ranged from .749 to .850, indicating adequate internal consistency for each construct. Discriminant validity was examined using the heterotrait-monotrait (HTMT) ratio (Henseler et al., 2015). Most HTMT ratios were below the .90 threshold recommended by Henseler et al. (2015). However, two pairs of constructs exceeded this threshold: AT-continuous intention (HTMT = .918) and responsibility-confidentiality (HTMT = .921). These elevated values indicate that discriminant validity was only partially established for these theoretically related construct pairs. The AT-continuous intention association is conceptually consistent with the TPB framework, in which AT functions as a proximal antecedent of intention. The responsibility-confidentiality association may reflect an overlap between ethical governance and data protection concerns in the use of GenAI. Accordingly,

Table 5. Hypothesis testing

#	Path	F ²	Estimate	t-value	p-value	Results
[H1]	AT → CI	0.427	0.557	21.688	***	Supported
[H2]	SN → CI	0.064	0.190	8.645	***	Supported
[H3]	PBC → CI	0.046	0.164	7.105	***	Supported
[H4]	A → AT	0.163	0.343	12.906	***	Supported
[H5]	O → AT	0.020	0.151	4.609	***	Supported
[H6]	R → AT	0.025	0.192	5.673	***	Supported
[H7]	C → AT	0.018	0.150	4.758	***	Supported

Note. CI: Continuous intention; A: Authenticity; O: Originality; R: Responsibility; C: Confidentiality; & ***p < .001

the structural paths involving these constructs should be interpreted with caution, and clearer boundaries between constructs should be examined in future research.

Model Fit

The fit indices for the structural model support an acceptable overall fit. The SRMR of 0.048 falls below the recommended threshold of 0.08, indicating an adequate approximate fit. The Chi-square statistic for the estimated model was 3710.626; given the well-documented sensitivity of this index to large samples, it is not treated as a primary fit criterion. These indices suggest that the structural model demonstrates an acceptable approximate fit and is suitable for hypothesis testing.

Hypothesis Testing

PLS-SEM was used in SmartPLS to test the seven hypotheses. **H1** was supported because AT positively predicted continuous intention to use GenAI ($\beta = .557$, $p < .001$). **H2** and **H3** were also supported: SN ($\beta = .190$, $p < .001$) and PBC ($\beta = .164$, $p < .001$) were significant positive predictors of continuous intention. Among the ethical antecedents of AT, **H4** was supported by the positive effect of authenticity ($\beta = .343$, $p < .001$), which showed the strongest effect among the ethical predictors. **H5**, **H6**, and **H7** were also supported, as originality ($\beta = .151$, $p < .001$), responsibility ($\beta = .192$, $p < .001$), and confidentiality ($\beta = .150$, $p < .001$) each positively predicted AT. The model explained 42.9% of the variance in AT and 63.8% of the variance in continuous intention. Effect sizes ranged from small to large, with the AT-to-continuous intention path showing the largest effect. **Table 5** shows the hypothesis testing results.

DISCUSSION AND CONCLUSIONS

Key Findings

GenAI tools have reshaped numerous fields by leveraging sophisticated language models to generate human-like content (Hwang et al., 2025). In academia and higher education, these AI systems have been the subject of recent research, and both their disadvantages and advantages have been widely discussed (Jin et al., 2025). This study investigated factors that influence the responsible implementation of AI in academia and higher education, with a particular focus on ethical considerations. Using a mixed-methods approach, this study first explored the advantages, limitations, and ethical considerations of GenAI use in academia and higher education through qualitative analysis.

Qualitative findings reveal that GenAI offers numerous advantages for researchers, including increased productivity, flexibility, and speed, as well as simplified work routines, text editing, translation, and language support. The technical limitations reported by participants highlight issues in cognitive processing, such as misinterpreting questions and decontextualization, as well as limitations in training and skills. The importance of ethical considerations, including originality, confidentiality, authenticity, and responsibility, was emphasized, underscoring the need for rigorous ethical standards to ensure responsible use in educational settings. These identified themes guided the development of the research model by revealing critical ethical factors.

In this study, TPB was chosen as the theoretical foundation, and the theory was extended to incorporate additional factors, such as ethical considerations. Quantitative results showed that AT toward the behavior,

SN, and PBC, which are the constructs of TPB, significantly predict individuals' continuous intention to use GenAI.

Perceptions of authenticity were positively associated with ATs toward AI use. Indeed, ChatGPT and other AI tools can generate false or deceptive information (Paul et al., 2023). Furthermore, AI tools can generate content that is not fact-based and sometimes provide unreliable details (Ariyaratne et al., 2023). Although AI tools can produce persuasive scientific summaries, they can also generate fabricated information and cite nonexistent or unpublished sources. This risk may be mitigated to some extent through customized prompts that guide the model more precisely and constrain the response context (Graf & Bernardi, 2023). These risks may reduce the perceived reliability of AI tools, suggesting that higher education institutions and researchers should adopt appropriate measures to limit misinformation (McDonald et al., 2025; Ryan, 2020).

Originality also emerged as a significant predictor of ATs toward GenAI use. Previous studies have consistently demonstrated that content generated by GenAI, without human verification, can contain significant errors and lack originality (Carvalho & Ivanov, 2024). This concern is consistent with evidence that GenAI tools may reproduce passages from online articles without proper source attribution (Alafnan et al., 2023).

Responsibility was associated with ATs toward GenAI use. Accordingly, rules and ethical standards should be developed, as interactions with GenAI may involve acquiring and storing sensitive data (Roberts et al., 2021). Furthermore, the use of GenAI raises questions of responsibility and accountability (McGreevey et al., 2020). The concept of responsibility emphasizes the need to ensure accountability as a prerequisite for addressing ethical concerns in the responsible integration of AI in academia and higher education (Floridi & Cows, 2022; Roberts et al., 2021).

Confidentiality further shaped users' ATs toward GenAI. There are growing concerns about security and privacy stemming from the possibility of misuse of data generated by these tools (Mohawesh et al., 2025). Previous studies on AI-enabled platforms have often emphasized the need to implement strong data protection measures to safeguard privacy (Adhikari et al., 2024; Sharma & Sharma, 2023).

Theoretical and Practical Implications

This study used a sequential mixed-methods approach in which structured interviews informed the development of an ethically grounded extension of the TPB model. The qualitative phase identified authenticity, originality, responsibility, and confidentiality as salient ethical concerns, which were then examined quantitatively as antecedents of AT. This design clarifies the study's specific contribution: it does not simply apply TPB to GenAI adoption but positions ethical concerns as measurable predictors within a continuous-use framework. The convergence between the qualitative themes and the quantitative model supports the relevance of these constructs, although the different participant profiles across phases require cautious interpretation.

The ethical constructs point to practical conditions under which GenAI can be integrated responsibly in higher education. Authenticity indicates the need for verification practices and attention to the reliability of AI-generated outputs (Hoffmann et al., 2024). In contrast, originality highlights the risk of unintentional plagiarism and the need for mechanisms that help users identify derivative content (Safi, 2025). Responsibility extends these concerns to institutional governance, user training, and clear rules for accountable use (Bin-Nashwan et al., 2025). Confidentiality introduces an additional requirement: students, researchers, and institutions should avoid uploading sensitive research data or intellectual property to GenAI systems unless appropriate privacy safeguards are in place (Arpaci, Aslan, et al., 2025). These findings suggest that higher education institutions should develop clear guidelines, provide training on ethical use, and involve relevant stakeholders in establishing policies for sustainable use of GenAI (Arpaci & Kusci, 2025).

This study addressed ethical concerns, including authenticity, responsibility, confidentiality, and originality. However, prior studies have also highlighted issues such as overreliance, loss of critical thinking skills, the digital divide, and technology dependency (Zhai et al., 2024). Overcoming all these ethical issues requires a multidisciplinary collaboration (Dwivedi et al., 2023). This collaboration aims to bridge the digital gap and ensure equal access to information for all, thereby promoting social equity and preserving pluralism (Farina et al., 2024). Therefore, higher education institutions should collaborate with relevant stakeholders to

establish clear codes of conduct and ethical guidelines governing the responsible and sustainable use of GenAI in research and higher education.

Limitations and Future Directions

This study has several limitations. First, because the data were self-reported, cross-sectional, and collected through convenience sampling, the findings should be interpreted as perceptual and associational rather than causal evidence. Second, the qualitative and quantitative phases involved different participant profiles: the interviews were conducted with faculty members, whereas the quantitative sample consisted predominantly of university students. This design is appropriate for instrument development, but it limits the direct transferability of qualitative themes to the student-based quantitative findings. The ethical constructs derived from faculty interviews should therefore be interpreted as a conceptual foundation for measurement rather than as evidence that faculty and students experience GenAI-related ethical issues in the same way. Future studies should test the model with matched or stratified samples across faculty, undergraduate, postgraduate, and institutional stakeholder groups.

The HTMT values for AT-continuous intention and responsibility-confidentiality exceeded the .90 threshold, indicating that discriminant validity was not fully established for these construct pairs. This limitation is meaningful because closely related constructs may share conceptual content, which can affect the interpretation of path coefficients. Future studies should refine the measurement items, test alternative validity criteria, and examine whether these constructs remain distinct across different higher education samples. Although Harman's single-factor test indicated that CMB was unlikely to pose a serious threat, this procedure is a limited diagnostic. CMB cannot be entirely ruled out, and future studies would benefit from additional procedural or statistical controls, such as marker-variable techniques or confirmatory approaches. Finally, alternative sampling strategies, including randomization or stratification, may enhance representativeness and improve the generalizability of the findings.

Author contributions: IA & MB: conceptualization, methodology, writing – original draft writing – review & editing. Both authors approved the final version of the article.

Funding: The authors received no financial support for the research and/or authorship of this article.

Ethics declaration: This study was approved by the Social and Human Sciences Ethics Committee at Bayburt University on 20 February 2025 with approval number 53. This study was conducted in accordance with the Declaration of Helsinki, and informed consent was obtained from all participants involved in the study.

AI statement: In this study, Grammarly Pro is used to improve readability and grammatical clarity with caution.

Declaration of interest: The authors declared no competing interest.

Data availability: Data generated or analyzed during this study are available from the authors on request.

REFERENCES

- Adhikari, K., Naik, N., Hameed, B. Z., Raghunath, S. K., & Somani, B. K. (2024). Exploring the ethical, legal, and social implications of ChatGPT in urology. *Current Urology Reports*, 25, 1-8. <https://doi.org/10.1007/s11934-023-01185-2>
- Ajzen, I. (1991). The theory of planned behavior. *Organizational Behavior and Human Decision Processes*, 50(2), 179-211. [https://doi.org/10.1016/0749-5978\(91\)90020-T](https://doi.org/10.1016/0749-5978(91)90020-T)
- Ajzen, I. (2002). Perceived behavioral control, self-efficacy, locus of control, and the theory of planned behavior. *Journal of Applied Social Psychology*, 32(4), 665-683. <https://doi.org/10.1111/j.1559-1816.2002.tb00236.x>
- Ajzen, I. (2012). The theory of planned behavior. In P. A. M. Van Lange, A. W. Kruglanski, & E. T. Higgins (Eds.), *Handbook of theories of social psychology* (vol. 1, pp. 438-459). SAGE. <https://doi.org/10.4135/9781446249215.n22>
- AlAfnan, M. A., Samira Dishari, Marina Jovic, & Koba Lomidze. (2023). ChatGPT as an educational tool: Opportunities, challenges, and recommendations for communication, business writing, and composition courses. *Journal of Artificial Intelligence and Technology*, 3(2), 60-68. <https://doi.org/10.37965/jait.2023.0184>

- Alalwan, A. A., Dwivedi, Y. K., Rana, N. P., & Algharabat, R. (2018). Examining factors influencing Jordanian customers' intentions and adoption of Internet banking: Extending UTAUT2 with risk. *Journal of Retailing and Consumer Services*, 40, 125-138. <https://doi.org/10.1016/j.jretconser.2017.08.026>
- Al-Emran, M., Al-Maroofof, R., Al-Sharafi, M. A., & Arpaci, I. (2022). What impacts learning with wearables? An integrated theoretical model. *Interactive Learning Environments*, 30(10), 1897-1917. <https://doi.org/10.1080/10494820.2020.1753216>
- Al-Emran, M., Al-Nuaimi, M. N., Arpaci, I., Al-Sharafi, M. A., & Anthony Jnr, B. (2023). Towards a wearable education: Understanding the determinants affecting students' adoption of wearable technologies using machine learning algorithms. *Education and Information Technologies*, 28(3), 2727-2746. <https://doi.org/10.1007/s10639-022-11294-z>
- Al-Emran, M., Arpaci, I., & Salloum, S. A. (2020). An empirical examination of continuous intention to use m-learning: An integrated model. *Education and Information Technologies*, 25(4), 2899-2918. <https://doi.org/10.1007/s10639-019-10094-2>
- Ali, I. M., & Arpaci, I. (2025). The role of moral intention and moral obligation in predicting attitudes toward avoiding Algiarism: A protection motivation theory perspective. *Computers in Human Behavior Reports*, 18, Article 100681. <https://doi.org/10.1016/j.chbr.2025.100681>
- Alkaissi, H., & McFarlane, S. I. (2023). *Artificial hallucinations in ChatGPT: Implications in scientific writing*. Cureus. <https://doi.org/10.7759/cureus.35179>
- Al-Maroofof, R. A., Arpaci, I., Al-Emran, M., Salloum, S. A., & Shaalan, K. (2021). Examining the acceptance of WhatsApp stickers through machine learning algorithms. *Studies in Systems, Decision and Control*, 295, 209-221. https://doi.org/10.1007/978-3-030-47411-9_12
- Al-Qaysi, N., Al-Emran, M., Al-Sharafi, M. A., Iranmanesh, M., Ahmad, A., & Mahmoud, M. A. (2025). Determinants of ChatGPT use and its impact on learning performance: An integrated model of BRT and TPB. *International Journal of Human-Computer Interaction*, 41(9), 5462-5474. <https://doi.org/10.1080/10447318.2024.2361210>
- Al-Sharafi, M. A., Al-Emran, M., Arpaci, I., Marques, G., Namoun, A., & Iahad, N. A. (2023). Examining the impact of psychological, social, and quality factors on the continuous intention to use virtual meeting platforms during and beyond COVID-19 pandemic: A hybrid SEM-ANN approach. *International Journal of Human-Computer Interaction*, 39(13), 2673-2685. <https://doi.org/10.1080/10447318.2022.2084036>
- Ariyaratne, S., Iyengar, K. P., Nischal, N., Chitti Babu, N., & Botchu, R. (2023). A comparison of ChatGPT-generated articles with human-written articles. *Skeletal Radiology*, 52(9), 1755-1758. <https://doi.org/10.1007/s00256-023-04340-5>
- Arpaci, I. (2019). A hybrid modeling approach for predicting the educational use of mobile cloud computing services in higher education. *Computers in Human Behavior*, 90, 181-187. <https://doi.org/10.1016/j.chb.2018.09.005>
- Arpaci, I. (2024). A multianalytical SEM-ANN approach to investigate the social sustainability of AI chatbots based on cybersecurity and protection motivation theory. *IEEE Transactions on Engineering Management*, 71, 1714-1725. <https://doi.org/10.1109/TEM.2023.3339578>
- Arpaci, I., & Kusci, I. (2025). A hybrid SEM-ANN approach for predicting the impact of psychological needs on satisfaction with generative AI use. *Technology, Knowledge and Learning*, 30, 2329-2350. <https://doi.org/10.1007/s10758-025-09817-x>
- Arpaci, I., & Sevinc, K. (2022). Development of the cybersecurity scale (CS-S): Evidence of validity and reliability. *Information Development*, 38(2), 218-226. <https://doi.org/10.1177/0266666921997512>
- Arpaci, I., Al-Emran, M., Al-Qaysi, N., & Al-Sharafi, M. A. (2025). What drives the use of generative artificial intelligence to promote educational sustainability? Evidence from SEM-ANN approach. *TechTrends*, 69(5), 957-971. <https://doi.org/10.1007/s11528-025-01089-7>
- Arpaci, I., Aslan, O., & Oner, I. E. (2025). Cybersecurity awareness scale (CSAS) for social media users: Development, validity and reliability study. *Information Development*. <https://doi.org/10.1177/02666669251336562>
- Atkins, S., Badrie, I., & Otterloo, S. (2021). Applying ethical AI frameworks in practice: Evaluating conversational AI chatbot solutions. *Computers and Society Research Journal*, 1. <https://doi.org/10.54822/QXOM4114>

- Belkina, M., Daniel, S., Nikolic, S., Haque, R., Lyden, S., Neal, P., Grundy, S., & Hassan, G. M. (2025). Implementing generative AI (GenAI) in higher education: A systematic review of case studies. *Computers and Education: Artificial Intelligence*, 8, Article 100407. <https://doi.org/10.1016/j.caeai.2025.100407>
- Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots. In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency* (pp. 610-623). ACM. <https://doi.org/10.1145/3442188.3445922>
- Bin-Nashwan, S. A., Sadallah, M., Benlahcene, A., & Ma'aji, M. M. (2025). ChatGPT and academic integrity: The role of personal best goals, academic competence, and workplace stress. *Foresight*, 27(4), 675-691. <https://doi.org/10.1108/FS-04-2024-0089>
- Bouschery, S. G., Blazevic, V., & Piller, F. T. (2023). Augmenting human innovation teams with artificial intelligence: Exploring transformer-based language models. *Journal of Product Innovation Management*, 40(2), 139-153. <https://doi.org/10.1111/jpim.12656>
- Carvalho, I., & Ivanov, S. (2024). ChatGPT for tourism: Applications, benefits and risks. *Tourism Review*, 79(2), 290-303. <https://doi.org/10.1108/TR-02-2023-0088>
- Chen, Y., Jensen, S., Albert, L. J., Gupta, S., & Lee, T. (2023). Artificial intelligence (AI) student assistants in the classroom: Designing chatbots to support student success. *Information Systems Frontiers*, 25(1), 161-182. <https://doi.org/10.1007/s10796-022-10291-4>
- Chiu, T. K. F., Xia, Q., Zhou, X., Chai, C. S., & Cheng, M. (2023). Systematic literature review on opportunities, challenges, and future research recommendations of artificial intelligence in education. *Computers and Education: Artificial Intelligence*, 4, Article 100118. <https://doi.org/10.1016/j.caeai.2022.100118>
- Chowdhury, M. M., Rifat, N., Ahsan, M., Latif, S., Gomes, R., & Rahman, M. S. (2023). ChatGPT: A threat against the CIA triad of cyber security. In *Proceedings of the 2023 IEEE International Conference on Electro Information Technology* (pp. 1-6). IEEE. <https://doi.org/10.1109/eIT57321.2023.10187355>
- Cotton, D. R. E., Cotton, P. A., & Shipway, J. R. (2023). Chatting and cheating: Ensuring academic integrity in the era of ChatGPT. *Innovations in Education and Teaching International*, 61(2), 228-239. <https://doi.org/10.1080/14703297.2023.2190148>
- Dave, T., Athaluri, S. A., & Singh, S. (2023). ChatGPT in medicine: An overview of its applications, advantages, limitations, future prospects, and ethical considerations. *Frontiers in Artificial Intelligence*, 6. <https://doi.org/10.3389/frai.2023.1169595>
- Deptula, A., Hunter, P. T., & Johnson-Sheehan, R. (2025). Rhetorics of authenticity: Ethics, ethos, and artificial intelligence. *Journal of Business and Technical Communication*, 39(1), 51-74. <https://doi.org/10.1177/10506519241280639>
- Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., Baabdullah, A. M., Koohang, A., Raghavan, V., Ahuja, M., Albanna, H., Albashrawi, M. A., Al-Busaidi, A. S., Balakrishnan, J., Barlette, Y., Basu, S., Bose, I., Brooks, L., Buhalis, D., ... Wright, R. (2023). "So what if ChatGPT wrote it?" Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *International Journal of Information Management*, 71, Article 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- Elkhatat, A. M. (2023). Evaluating the authenticity of ChatGPT responses: A study on text-matching capabilities. *International Journal for Educational Integrity*, 19, Article 15. <https://doi.org/10.1007/s40979-023-00137-0>
- English, R., Nash, R., & Mackenzie, H. (2025). 'A rather stupid but always available brain-storming partner': Use and understanding of generative AI by UK postgraduate researchers. *Innovations in Education and Teaching International*, 63(1), 193-207. <https://doi.org/10.1080/14703297.2024.2446236>
- Farina, M., Yu, X., & Lavazza, A. (2024). Ethical considerations and policy interventions concerning the impact of generative AI tools in the economy and in society. *AI and Ethics*, 5, 737-745. <https://doi.org/10.1007/S43681-023-00405-2>
- Farrokhnia, M., Banihashem, S. K., Noroozi, O., & Wals, A. (2023). A SWOT analysis of ChatGPT: Implications for educational practice and research. *Innovations in Education and Teaching International*, 61(3), 460-474. <https://doi.org/10.1080/14703297.2023.2195846>
- Floridi, L. and Cows, J. (2022). A unified framework of five principles for AI in society. In S. Carta (Ed.), *Machine learning and the city*. Wiley. <https://doi.org/10.1002/9781119815075.ch45>

- Gaffar, H., & Albarashdi, S. (2024). Copyright protection for AI-generated works: Exploring originality and ownership in a digital landscape. *Asian Journal of International Law*, 15(1), 23-46. <https://doi.org/10.1017/S2044251323000735>
- Ganguly, A., Johri, A., Ali, A., & McDonald, N. (2025). Generative artificial intelligence for academic research: Evidence from guidance issued for researchers by higher education institutions in the United States. *AI and Ethics*, 5, 3917-3933. <https://doi.org/10.1007/S43681-025-00688-7>
- Gašević, D., Siemens, G., & Sadiq, S. (2023). Empowering learners for the age of artificial intelligence. *Computers and Education: Artificial Intelligence*, 4, Article 100130. <https://doi.org/10.1016/j.caeai.2023.100130>
- Gibreel, O., & Arpacı, I. (2025). Development and validation of the prompt engineering competence scale (PECS). *Information Development*. <https://doi.org/10.1177/02666669251336455>
- Graf, A., & Bernardi, R. E. (2023). ChatGPT in research: Balancing ethics, transparency and advancement. *Neuroscience*, 515, 71-73. <https://doi.org/10.1016/j.neuroscience.2023.02.008>
- Hair, J. F., Black, W. C., Babin, B. J., & Anderson, R. E. (2019). *Multivariate data analysis* (8th ed.). Cengage Learning.
- Henseler, J., Ringle, C. M., & Sarstedt, M. (2015). A new criterion for assessing discriminant validity in variance-based structural equation modeling. *Journal of the Academy of Marketing Science*, 43(1), 115-135. <https://doi.org/10.1007/s11747-014-0403-8>
- Hill, J., Randolph Ford, W., & Farreras, I. G. (2015). Real conversations with artificial intelligence: A comparison between human-human online conversations and human-chatbot conversations. *Computers in Human Behavior*, 49, 245-250. <https://doi.org/10.1016/j.chb.2015.02.026>
- Hoffmann, S., Lasarov, W., & Dwivedi, Y. K. (2024). AI-empowered scale development: Testing the potential of ChatGPT. *Technological Forecasting and Social Change*, 205, Article 123488. <https://doi.org/10.1016/j.techfore.2024.123488>
- Huang, W., Hew, K. F., & Fryer, L. K. (2022). Chatbots for language learning—Are they really useful? A systematic review of chatbot-supported language learning. *Journal of Computer Assisted Learning*, 38(1), 237-257. <https://doi.org/10.1111/jcal.12610>
- Hutson, M. (2021). Robo-writers: The rise and risks of language-generating AI. *Nature*, 591(7848), 22-25. <https://doi.org/10.1038/d41586-021-00530-0>
- Hwang, Y., & Jeong, S. H. (2025). Generative artificial intelligence and misinformation acceptance: An experimental test of the effect of forewarning about artificial intelligence hallucination. *Cyberpsychology, Behavior, and Social Networking*, 28(4), 284-289. <https://doi.org/10.1089/cyber.2024.0407>
- Hwang, Y., Shin, D., & Lee, J. H. (2025). Who owns AI-generated artwork? Revisiting the work of generative AI based on human-AI cocreation. *Telematics and Informatics*, 98, Article 102266. <https://doi.org/10.1016/j.tele.2025.102266>
- Janssen, M. (2025). Responsible governance of generative AI: Conceptualizing GenAI as complex adaptive systems. *Policy and Society*, 44(1), 38-51. <https://doi.org/10.1093/polsoc/puae040>
- Jin, Y., Yan, L., Echeverria, V., Gašević, D., & Martinez-Maldonado, R. (2025). Generative AI in higher education: A global perspective of institutional adoption policies and guidelines. *Computers and Education: Artificial Intelligence*, 8, Article 100348. <https://doi.org/10.1016/j.caeai.2024.100348>
- Johnson, D. G., & Noorman, M. (2014). Recommendations for future development of artificial agents. *IEEE Technology and Society Magazine*, 33(4), 22-28. <https://doi.org/10.1109/MTS.2014.2363978>
- Kasneci, E., Sessler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., Gasser, U., Groh, G., Günnemann, S., Hüllermeier, E., Krusche, S., Kutyniok, G., Michaeli, T., Nerdel, C., Pfeffer, J., Poquet, O., Sailer, M., Schmidt, A., Seidel, T., ... Kasneci, G. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103, Article 102274. <https://doi.org/10.1016/j.lindif.2023.102274>
- Khalil, M., & Er, E. (2023). Will ChatGPT get you caught? Rethinking of plagiarism detection. In P. Zaphiris, & A. Ioannou (Eds.), *Learning and collaboration technologies. HCL 2023. Lecture notes in computer science, vol 14040* (pp. 475-487). Springer. https://doi.org/10.1007/978-3-031-34411-4_32
- Kim, J., Yu, S., Detrick, R., & Li, N. (2024). Exploring students' perspectives on generative AI-assisted academic writing. *Education and Information Technologies*, 30, 1265-1300. <https://doi.org/10.1007/s10639-024-12878-7>

- King, M. R. (2023). A conversation on artificial intelligence, chatbots, and plagiarism in higher education. *Cellular and Molecular Bioengineering*, 16, 1-2. <https://doi.org/10.1007/s12195-022-00754-8>
- Kohnke, L., Moorhouse, B. L., & Zou, D. (2023). ChatGPT for language teaching and learning. *RELC Journal*, 54(2), 537-550. <https://doi.org/10.1177/00336882231162868>
- Kuhail, M. A., Alturki, N., Alramlawi, S., & Alhejori, K. (2023). Interacting with educational chatbots: A systematic review. *Education and Information Technologies*, 28(1), 973-1018. <https://doi.org/10.1007/s10639-022-11177-3>
- Kumar, S., Rao, P., Singhania, S., Verma, S., & Kheterpal, M. (2024). Will artificial intelligence drive the advancements in higher education? A tri-phased exploration. *Technological Forecasting and Social Change*, 201, Article 123258. <https://doi.org/10.1016/j.TECHFORE.2024.123258>
- Lund, B. D., & Wang, T. (2023). Chatting about ChatGPT: How may AI and GPT impact academia and libraries? *Library Hi Tech News*, 40(3), 26-29. <https://doi.org/10.1108/LHTN-01-2023-0009>
- Maras, M.-H., & Alexandrou, A. (2019). Determining authenticity of video evidence in the age of artificial intelligence and in the wake of deepfake videos. *The International Journal of Evidence & Proof*, 23(3), 255-262. <https://doi.org/10.1177/1365712718807226>
- McDonald, N., Johri, A., Ali, A., & Collier, A. H. (2025). Generative artificial intelligence in higher education: Evidence from an analysis of institutional policies and guidelines. *Computers in Human Behavior: Artificial Humans*, 3, Article 100121. <https://doi.org/10.1016/J.CHBAH.2025.100121>
- McGreevey, J. D., Hanson, C. W., & Koppel, R. (2020). Clinical, legal, and ethical aspects of artificial intelligence-assisted conversational agents in health care. *JAMA*, 324(6), Article 552. <https://doi.org/10.1001/jama.2020.2724>
- McTear, M. (2021). *Conversational AI*. Springer. <https://doi.org/10.1007/978-3-031-02176-3>
- Mhlanga, D. (2023). *Open AI in education, the responsible and ethical use of ChatGPT towards lifelong learning*. SSRN. <https://doi.org/10.2139/ssrn.4354422>
- Mohawesh, R., Ottom, M. A., & Salameh, H. B. (2025). A data-driven risk assessment of cybersecurity challenges posed by generative AI. *Decision Analytics Journal*, 15, Article 100580. <https://doi.org/10.1016/J.DAJOUR.2025.100580>
- Nguyen, K. V. (2025). The use of generative AI tools in higher education: Ethical and pedagogical principles. *Journal of Academic Ethics*, 23, 1435-1455. <https://doi.org/10.1007/S10805-025-09607-1>
- Nguyen, T. H., Waizenegger, L., & Techatassanasoontorn, A. A. (2022). "Don't neglect the user!" – Identifying types of human-chatbot interactions and their associated characteristics. *Information Systems Frontiers*, 24(3), 797-838. <https://doi.org/10.1007/s10796-021-10212-x>
- Paul, J., Ueno, A., & Dennis, C. (2023). ChatGPT and consumers: Benefits, pitfalls and future research agenda. *International Journal of Consumer Studies*, 47(4), 1213-1225. <https://doi.org/10.1111/ijcs.12928>
- Podsakoff, P. M., MacKenzie, S. B., & Podsakoff, N. P. (2012). Sources of method bias in social science research and recommendations on how to control it. *Annual Review of Psychology*, 63(1), 539-569. <https://doi.org/10.1146/annurev-psych-120710-100452>
- Prinz, K. D. (2025). Managing the legal risks of artificial intelligence on intellectual property and confidential information. *Consulting Psychology Journal*, 77(2), 169-179. <https://doi.org/10.1037/cpb0000287>
- Qin, F., Li, K., & Yan, J. (2020). Understanding user trust in artificial intelligence-based educational systems: Evidence from China. *British Journal of Educational Technology*, 51(5), 1693-1710. <https://doi.org/10.1111/bjet.12994>
- Roberts, H., Cows, J., Morley, J., Taddeo, M., Wang, V., & Floridi, L. (2021). The Chinese approach to artificial intelligence: An analysis of policy, ethics, and regulation. *AI & SOCIETY*, 36(1), 59-77. <https://doi.org/10.1007/s00146-020-00992-2>
- Ryan, M. (2020). In AI we trust: Ethics, artificial intelligence, and reliability. *Science and Engineering Ethics*, 26(5), 2749-2767. <https://doi.org/10.1007/s11948-020-00228-y>
- Safi, R. (2025). Detecting plagiarism in the age of generative AI: An exploratory experiment. *Communications of the Association for Information Systems*, 56(1), 594-612. <https://doi.org/10.17705/1CAIS.05624>
- Sharma, M., & Sharma, S. (2023). Transforming maritime health with ChatGPT-powered healthcare services for mariners. *Annals of Biomedical Engineering*, 51(6), 1123-1125. <https://doi.org/10.1007/s10439-023-03195-0>

- Shen, Y., Heacock, L., Elias, J., Hentel, K. D., Reig, B., Shih, G., & Moy, L. (2023). ChatGPT and other large language models are double-edged swords. *Radiology*, 307(2), Article e230163. <https://doi.org/10.1148/radiol.230163>
- Söllner, M., Arnold, T., Benlian, A., Bretschneider, U., Knight, C., Ohly, S., Rudkowski, L., Schreiber, G., & Wendt, D. (2025). ChatGPT and beyond: Exploring the responsible use of generative AI in the workplace: An interdisciplinary perspective. *Business and Information Systems Engineering*, 67(2), 289-303. <https://doi.org/10.1007/s12599-025-00932-8>
- Stokel-Walker, C., & Van Noorden, R. (2023). What ChatGPT and generative AI mean for science. *Nature*, 614(7947), 214-216. <https://doi.org/10.1038/d41586-023-00340-6>
- Storey, V. C. (2024). Knowledge management in a world of generative AI: Impact and implications. *ACM Transactions on Management Information Systems*, 16(3), Article 26. <https://doi.org/10.1145/3719209>
- Tabachnick, B. G., & Fidell, L. S. (2012). *Using multivariate statistics* (6th ed.). Pearson.
- Ulnicane, I. (2025). Governance fix? Power and politics in controversies about governing generative AI. *Policy and Society*, 44(1), 70-84. <https://doi.org/10.1093/POLSOC/PUAE022>
- Vorobyeva, K. I., Belous, S., Savchenko, N. V., Smirnova, L. M., Nikitina, S. A., & Zhdanov, S. P. (2025). Personalized learning through AI: Pedagogical approaches and critical insights. *Contemporary Educational Technology*, 17(2), Article ep574. <https://doi.org/10.30935/CEDTECH/16108>
- Wang, C., Liu, S., Yang, H., Guo, J., Wu, Y., & Liu, J. (2023). Ethical considerations of using ChatGPT in health care. *Journal of Medical Internet Research*, 25, Article e48009. <https://doi.org/10.2196/48009>
- Wang, J., & Fan, W. (2025). The effect of ChatGPT on students' learning performance, learning perception, and higher-order thinking: Insights from a meta-analysis. *Humanities and Social Sciences Communications*, 12, Article 621. <https://doi.org/10.1057/s41599-025-04787-y>
- Wu, H., & Dong, Z. (2025). What motivates second language majors to use generative AI for informal learning? Insights from the theory of planned behavior. *IEEE Access*, 13, 34877-34886. <https://doi.org/10.1109/ACCESS.2025.3544489>
- Yeh, H.-C. (2025). The synergy of generative AI and inquiry-based learning: Transforming the landscape of English teaching and learning. *Interactive Learning Environments*, 33(1), 88-102. <https://doi.org/10.1080/10494820.2024.2335491>
- Zhai, C., Wibowo, S., & Li, L. D. (2024). The effects of over-reliance on AI dialogue systems on students' cognitive abilities: A systematic review. *Smart Learning Environments*, 11, Article 28. <https://doi.org/10.1186/s40561-024-00316-7>
- Zhou, J., Ke, P., Qiu, X., Huang, M., & Zhang, J. (2024). ChatGPT: Potential, prospects, and limitations. *Frontiers of Information Technology & Electronic Engineering*, 25(1), 6-11. <https://doi.org/10.1631/fitee.2300089>
- Zhou, J., Müller, H., Holzinger, A., & Chen, F. (2023). Ethical ChatGPT: Concerns, challenges, and commandments. *Electronics*, 13(17), Article 3417. <https://doi.org/10.3390/electronics13173417>

APPENDIX A: SCALE ITEMS

Authenticity

- A1. "GenAI tools should generate information with high accuracy and credibility."
- A2. "I am concerned about the accuracy of the information provided by GenAI tools."
- A3. "GenAI tools should consistently generate reliable and trustworthy information."

Originality

- O1. "I am concerned that GenAI tools may inadvertently generate content that resembles or mirrors existing copyrighted material."
- O2. "GenAI tools should have mechanisms to detect and prevent potential plagiarism in the content they generate."
- O3. "I worry about the possibility of unintentional plagiarism when using GenAI tools."

Responsibility

- R1. "I am concerned about the potential negative impact GenAI tools might have on academic work if used carelessly."
- R2. "GenAI tools should come with guidelines to prevent misuse for unethical purposes."
- R3. "It is essential to educate users about the possible negative impacts and how to avoid misuse when utilizing GenAI tools."
- R4. "It is important to me that the use of GenAI tools aligns with academic integrity and ethical standards."

Confidentiality

- C1. "The security of my data and interactions while using GenAI tools is a significant concern for me."
- C2. "GenAI tools should implement robust security measures to protect user data and privacy."
- C3. "I worry about potential breaches of privacy or unauthorized access to data when using GenAI tools."

Attitude

- AT1. "Using GenAI tools is a wise decision."
- AT2. "I view using GenAI tools favorably."
- AT3. "Using GenAI tools would be enjoyable."
- AT4. "Learning to use GenAI tools can be a fulfilling experience."

Subjective Norms

- SN1. "People who are important to me think that I should use GenAI tools."
- SN2. "People who influence my behavior think that I should use GenAI tools."
- SN3. "People whose opinions I value prefer that I use GenAI tools."

Perceived Behavioral Control

- PBC1. "I have sufficient knowledge to use GenAI tools."
- PBC2. "I have sufficient control to decide to use GenAI tools."
- PBC3. "I have sufficient self-confidence to decide to use GenAI tools."

Continuous Intention

- CI1. "I am likely to continue using GenAI tools in the future."
- CI2. "I intend to continue using GenAI tools in my daily life."
- CI3. "I have a strong desire to continue using GenAI tools in the future."

