

Research Article:

The Ethical Use of AI Technology Applications Among University Students in Pahang

Mohd Rashimi Mohamed^{1*}, Norazmi Anas^{2*}, Zuriani Yaacob³, Ahmad Hafizi Ahmad Giran⁴, Norazman Lot⁵, Ahmad Effat Mokhtar⁶ and Zulkifli Mohd Ghazali⁷

¹Faculty of Quranic Science, UCYP University, 26060 Kuantan, Pahang, Malaysia

²Academy of Contemporary Islamic Studies, Universiti Teknologi MARA, Perak Branch, Tapah Campus, 35400 Tapah Road, Perak, Malaysia

³Akademi Pengajian Bahasa, Universiti Teknologi MARA, Pahang Branch, Raub Campus, 27600 Raub, Pahang, Malaysia

⁴Department of Business and Accountancy, Kolej Poly-Tech MARA Kuantan, 25150 Kuantan, Pahang, Malaysia

⁵Department of General Studies, Faculty of Education and General Studies, UCYP University, 26060 Kuantan, Pahang, Malaysia

⁶Department of Islamic Studies, Universiti Al-Quran Al-Sultan Abdullah Ahmad Shah Pahang, 25150 Kuantan, Pahang, Malaysia

⁷Faculty of Computer and Mathematical Sciences, Universiti Teknologi MARA, Perak Branch, Tapah Campus, 35400 Tapah Road, Perak, Malaysia

*Corresponding authors: mohd_rashimi@ucyp.edu.my; norazmianas@uitm.edu.my

ABSTRACT

Artificial Intelligence (AI) tools are increasingly integrated into higher education, raising questions about how ethically students engage with such technologies. This study examines the factors that influence ethical use of AI among university students in Pahang, Malaysia. A total of 467 students from eight institutions participated in an online survey, and the responses were analysed using Partial Least Squares Structural Equation Modelling (PLS-SEM). Four key constructs were evaluated: perceived risk, verification intention, negative emotion, and responsible AI use. The results indicate that a higher perceived risk is associated with a stronger intention to verify AI-generated content, and that this verification intention significantly predicts responsible AI use. Negative emotions also positively influence responsible use and moderate the relationship between verification intention and responsible AI behaviour. These findings suggest that both cognitive assessments and emotional responses play significant roles in shaping how students manage AI-generated outputs and avoid unethical academic practices. This study contributes to current discourse by shifting attention from AI adoption to post-adoption ethical behaviour, demonstrating how risk perception and emotional awareness encourage verification practices and responsible engagement. The findings offer practical insights for universities aiming to promote ethical AI literacy and foster academic integrity in digital learning environments.

Keywords: Ethical use, AI technology, students, higher education, state of Pahang

Published: 9 June 2026

To cite this article: Mohamed, M. R., Anas, N., Yaacob, Z., Ahmad Giran, A. H., Lot, N., Mokhtar, A. E., & Ghazali, Z. M. (2026). The ethical use of AI technology applications among university students in Pahang. *Asia Pacific Journal of Educators and Education*, 41(1), 275–296. <https://doi.org/10.21315/apjee2026.41.1.12>

INTRODUCTION

Artificial intelligence (AI) has rapidly transformed teaching and learning practices in higher education. The emergence of generative AI tools such as ChatGPT has significantly influenced how students access information, complete assignments, and engage with academic materials. These tools support idea generation, facilitate understanding of complex concepts and provide writing assistance, thereby enhancing learning efficiency and accessibility (Kasneci et al., 2023; Cotton et al., 2023).

Despite these benefits, the increasing use of AI tools has raised ethical concerns in academic contexts. Scholars have highlighted risks related to misinformation, plagiarism, algorithmic bias, and excessive reliance on AI-generated outputs (Dwivedi et al., 2023; Susnjak & McIntosh, 2024). AI systems may produce inaccurate or fabricated information if their outputs are accepted without verification (Li et al., 2023). These issues have intensified debates about academic integrity, responsible use of technology, and the ethical boundaries of AI-assisted learning in higher education. Although previous research has widely examined students' adoption and perceptions of AI technologies, comparatively less attention has been given to students' ethical behaviour after adopting these tools. In particular, research on how students verify AI-generated information and regulate their use of AI in academic tasks remains limited. Understanding these post-adoption behaviours is essential for promoting responsible AI engagement in educational environments.

In Malaysia, the growing accessibility of generative AI tools has similarly influenced students' academic practices. AI applications are increasingly used to assist with writing, idea generation, and academic research tasks. However, concerns have also emerged regarding the misuse of AI tools to complete academic assignments and their potential implications for academic integrity (Mustapha et al., 2024; Maulana et al., 2023). Considering these concerns, this study investigates the ethical use of AI applications among university students in Pahang, Malaysia. To address this need, the present study pursues the following objectives:

RO1: Identify the most used AI technology applications among university students in the state of Pahang.

RO2: Examine the ethical use of AI technology applications among university students in the state of Pahang.

LITERATURE REVIEW

Integration of Artificial Intelligence in Higher Education

Commonly, artificial intelligence (AI) applications cover ChatGPT, Grammarly, Quillbot, and Socratic (Selvi et al., 2024). The rapid growth of AI has helped students in higher

education, driving broad changes in many countries. For example, in Canada, Martiniello et al. (2021) examined AI tools to support students with mobility disabilities, whereas in Azerbaijan, Imran et al. (2024) took a different approach. AI also supports student mental health and career advice. Australia uses predictive analytics, focusing on supervised learning to track and improve student performance (Fahd et al., 2022; Bond et al., 2024). In the Gulf Cooperation Council (GCC) region, researchers found advance machine learning methods such as Random Forests and Support Vector Machines to predict student outcomes (Fadlelmula & Qadhi, 2024).

Other regions have focused on pedagogical dimensions of AI adoption. In China, Long et al. (2025) explain how AI enhances student engagement through flipped classrooms and project-based learning. country of Austria highlights the institutional infrastructure underpinning AI integration, particularly through learning management systems (Kuka et al., 2022). In Argentina, Ramirez and Esparrell (2024) highlighted the use of multi-agent systems to personalise learning pathways. European scholarship study conducted by Marengo et al. (2024) emphasises the use of chatbots to enhance students' engagement with a validated AI tool. In the Southeast Asian context, Zhao et al. (2025) provide valuable insights into a cross-country comparative analysis of Malaysia and China, revealing notable differences. Malaysian students use AI to improve English language skills and structure assignments, while Chinese students have far more diverse applications, focusing on programming support, artistic projects, and simulation exercises. Collectively, AI integration varies across regions with a united mission of improving students' experience in learning for student outcomes.

Ethical Challenges of AI Use in Educational Contexts

AI poses genuine challenges to academic integrity in diverse areas. Initially, the lens focuses on academic dishonesty and cheating behaviour. Several studies literally stated that students purposely used generative AI to produce written work dishonestly. Hanh and Duyen (2025) identified two distinct forms of AI-assisted cheating: misusing AI to complete academic tasks and improperly incorporating AI-generated content. The ways of dishonesty make essay-writing assignments easier, where chatbots like ChatGPT can generate entire papers (King & ChatGPT, 2023). Notably, this is how cultural challenges in education focus. Students view AI integration for writing entire papers as major misconduct, and it's quite surprising that they see smaller AI-assisted tasks as less severe (Lund et al., 2025). Different findings called, 50% of students believe extensive AI use is unethical, even as nearly 20% engage in it (Garrote Jurado et al., 2023). In the worst-case scenario, students noted that misconduct involving generative AI, and 50% of them were less likely to view it as cheating (Tindle et al., 2023).

Beyond individual misconduct, previous studies centre on broader systemic challenges. Elmessiry et al. (2023) raise concerns about algorithmic bias. Dong et al. (2025) take a different view of learning burnout and information overload. Hanh and Duyen (2025) attribute the structural factors to work pressure and ethical ambiguity. The unethical

culture highlighted needs to be addressed, not resting solely on the student's shoulders as a moral failure. Systemic conditions also need to be examined to build students' awareness of this unethical use of AI.

Students' Behavioural and Psychological Responses to AI Use

Students' interactions with AI technologies are influenced not only by technological factors but also by cognitive and psychological processes. Research suggests that students' perceptions of risk, trust, and usefulness significantly shape their willingness to engage with emerging technologies (Zawacki-Richter et al., 2019). As generative AI tools become more widely accessible, understanding how students perceive and regulate their use of these technologies has become increasingly important. Perceived risk represents an important cognitive factor influencing responsible technology use. When students recognise potential risks such as misinformation, hallucinated outputs, or biased recommendations, they may adopt more cautious behaviours when interacting with AI-generated content (Li et al., 2023; Kasneci et al., 2023). This awareness may encourage verification intention, which refers to the willingness to critically evaluate and confirm the accuracy of AI-generated information before applying it in academic work (Susnjak & McIntosh, 2024). Verification practices play a critical role in ensuring that AI tools are used as supportive learning resources rather than substitutes for intellectual effort.

Emotional responses may also influence ethical behaviour in digital environments. Feelings of concern, discomfort, or anxiety regarding AI misuse may prompt students to reflect more carefully on their use of AI technologies and encourage more responsible decision-making (Chen & Lin, 2024). These emotional responses may interact with cognitive evaluations of risk, shaping how students regulate their behaviour when engaging with AI tools. Responsible AI use in academic contexts therefore involves behaviours aligned with academic integrity, including verifying information, acknowledging AI assistance, and ensuring that AI-generated outputs complement rather than replace independent learning (Kasneci et al., 2023). Understanding the psychological and behavioural mechanisms underlying these practices is essential for promoting ethical engagement with AI technologies in educational settings.

Theoretical Perspective and Research Gap

The Technology Acceptance Model (TAM) provides an established framework for understanding how individuals adopt and use new technologies. Traditionally, TAM emphasises perceived usefulness and perceived ease of use as key determinants of technology adoption (Davis, 1989). However, as AI technologies become widely accessible, the research focus has increasingly shifted beyond adoption to examine how users regulate their behaviour after adopting these tools. Recent discussions suggest that the study of AI in education should move beyond simple adoption models to examine post-adoption behaviours related to responsibility, critical evaluation, and ethical decision-making (Dwivedi et al., 2023). From an ethical standpoint, AI technologies operate within broader

socio-technical systems in which users remain accountable for their decisions and actions (Floridi, 2018). Ethical AI frameworks, therefore, emphasise principles such as fairness, transparency, accountability, and responsible human oversight (Floridi & Cowls, 2022).

Despite increasing scholarly attention to AI integration in education, much of the existing literature has focused on students' attitudes toward AI, adoption intentions, and perceived usefulness. Compared with other studies, fewer have examined how cognitive and emotional mechanisms influence students' ethical behaviour when using AI tools in academic settings. In particular, limited research has investigated how perceived risk, verification intention, and emotional responses jointly shape responsible AI usage among university students. To address this gap, the present study investigates the ethical use of AI technology applications among university students in Pahang, Malaysia, focusing on the role of perceived risk, verification intention, and emotional responses in shaping responsible AI behaviour.

METHODS

This study adopts a quantitative case study design to examine the ethical use of AI technology applications among selected university students in the state of Pahang. The primary objective is to gather measurable data regarding students' levels of awareness, understanding, and application of ethical principles when interacting with AI tools, particularly in academic contexts. By focusing on a specific population in Pahang, the study aims to identify patterns, trends, and potential gaps in ethical practices, which may reflect broader issues in digital literacy and responsible technology use. The case study approach enables a focused examination of real-world behaviours. At the same time, the quantitative methodology ensures that the findings are statistically valid and can inform future policy development, curriculum enhancement, and institutional guidelines related to AI ethics in higher education.

Theoretical Foundation

The TAM provides one of the most widely used theoretical frameworks for explaining how individuals adopt and use new technologies. According to TAM, users' behavioural intention to use technology is influenced primarily by their perceptions of usefulness and ease of use, which subsequently determine actual usage behaviour (Davis, 1989). Due to its strong explanatory power, TAM has been widely applied to examine the adoption of digital technologies in various contexts, including educational environments.

However, the rapid proliferation of generative AI tools has shifted the focus of research beyond initial technology adoption. Applications such as ChatGPT, Grammarly, and other AI-assisted learning platforms are now widely accessible and frequently used by university students for academic activities, including writing support, idea generation, and information search. Consequently, the key issue in higher education is no longer whether students adopt AI technologies, but rather how they engage with these tools responsibly

within academic settings. Recent discussions highlight that research on generative AI should therefore extend beyond adoption models to examine how users regulate their behaviour after adoption, particularly in relation to responsible and ethical use of AI technologies in educational contexts (Dwivedi et al., 2023; Kasneci et al., 2023).

To address this emerging concern, the present study adopts a post-adoption behavioural perspective grounded in the behavioural logic of TAM. While TAM traditionally focuses on factors influencing technology adoption, its core principle that behavioural intention predicts actual behaviour remains relevant in post-adoption contexts. In this study, verification intention represents students' behavioural intention to critically evaluate AI-generated information before using it in academic tasks, whereas responsible use of AI represents the behavioural outcome reflecting ethical and responsible engagement with AI technologies.

Building on this behavioural pathway, the study incorporates additional cognitive and emotional factors that may influence responsible AI usage. Perceived risk reflects students' awareness of potential problems associated with AI-generated outputs, such as misinformation, bias, or academic misuse. Students who perceive greater risks when using AI technologies may become more cautious and develop stronger intentions to verify AI-generated information before applying it in academic work. In addition to cognitive evaluation, negative emotions, such as concern or discomfort about AI misuse, may influence ethical decision-making by encouraging more careful, reflective behaviour when interacting with AI tools.

Together, these constructs explain how students regulate their behaviour when engaging with AI technologies in academic contexts. Rather than focusing solely on technology adoption, the proposed model emphasises behavioural mechanisms that promote responsible use of AI. By integrating perceived risk, verification intention, emotional responses, and responsible AI behaviour, this study extends the TAM perspective toward understanding ethical engagement with AI technologies in higher education.

Sample and Data Collection

This study examined university students' perceptions of the ethical use of AI technologies in academic contexts. The target population consisted of students enrolled in higher education institutions located in the state of Pahang, Malaysia. University students were selected as the study population because they represent one of the most active user groups of AI-assisted learning tools such as writing assistants, productivity applications, and AI-based information generators. The study employed a cross-sectional survey design using an online questionnaire. Data were collected from students across eight higher education institutions in Pahang, including University College of Yayasan Pahang (UCYP), Universiti Islam Pahang Sultan Ahmad Shah (UnIPSAS), Universiti Teknologi MARA (UiTM), International Islamic University Malaysia (IIUM), Universiti Malaysia Pahang Al-Sultan Abdullah (UMPSA), Kolej Poly-Tech MARA (KPTM), Kolej Yayasan Pahang (KYP), and Kolej Kemahiran Tinggi MARA (KKTM).

The survey link was distributed through institutional communication channels and student associations, enabling students from various academic programmes and study levels to voluntarily participate. This approach enabled broad access to the target population and facilitated participation from multiple institutions within the region. Participation was voluntary and anonymous, and informed consent was obtained electronically before respondents proceeded with the questionnaire. A total of 467 valid responses were obtained and included in the final analysis. Prior to analysis, responses were screened to remove incomplete submissions and patterned responses to ensure data quality. Although the sampling approach enabled broad participation across institutions, it may introduce sampling bias due to the voluntary nature of survey participation. Therefore, the findings should be interpreted within the context of the study population, which primarily represents students enrolled in higher education institutions within the state of Pahang.

Measures

The measures employed in the questionnaire were adapted from established research frameworks in educational technology and digital ethics to ensure both theoretical alignment and content validity (Creswell & Creswell, 2018). To strengthen measurement clarity, three academic experts with backgrounds in educational technology, ethics, and quantitative methodology reviewed the instrument to evaluate item relevance, language accuracy, and conceptual coverage. Minor wording refinements were made based on their feedback to improve clarity and contextual appropriateness.

The survey instrument consisted of demographic items that captured age, gender, academic background, and level of study, providing contextual insight into participants' profiles. It further included items measuring the main constructs related to ethical AI usage, namely Perceived Risk, Verification Intention, Negative Emotion, and Responsible Use of AI. The final component comprised items appropriate for psychometric assessment and subsequent modelling using Partial Least Squares Structural Equation Modelling (PLS-SEM). All constructed items were assessed using a 5-point Likert scale ranging from 1 (Strongly Disagree) to 5 (Strongly Agree), consistent with best practices for attitudinal and behavioural measurement.

Following data collection, initial descriptive analysis and reliability screening were conducted using the Statistical Package for the Social Sciences (SPSS) version 26 (Pallant, 2020). Subsequently, PLS-SEM was applied to evaluate both the measurement and structural models. This analytical approach is suitable for exploratory research involving multiple latent constructs with reflective indicators. The analysis was conducted using SmartPLS 4.0, following the recommended procedures (Hair et al., 2019), which included examining outer loadings, internal consistency reliability (Cronbach's alpha and composite reliability), convergent validity (Average Variance Extracted, AVE), and discriminant validity based on the Fornell–Larcker criterion. This combined strategy ensured robust assessment of the instrument's psychometric soundness and supported the evaluation of theoretical relationships among constructs within the context of ethical AI usage among university students.

RESULTS

This section presents the study’s findings based on data collected from 467 university students regarding their ethical use of AI technologies in academic settings. The results are organised to address the stated research objectives and hypotheses. First, the demographic characteristics of the respondents are summarised to provide context for the sample population. This is followed by a descriptive analysis of the four key constructs examined, with the results presented in Tables 2 to 5. The measurement model is then assessed to establish the reliability and validity of the constructs. Subsequently, the structural model is evaluated to examine the relationships among the constructs. Finally, the discussion section interprets these findings in relation to existing theories and literature, highlighting their theoretical and practical implications for promoting the ethical use of AI among students.

Demographic Information

Table 1 presents the demographic profile of the respondents. The sample comprises 467 university students from eight higher education institutions in Pahang, with a diverse distribution across gender, age group, programme level, and institution. The results also indicate that students actively use various AI tools in their academic activities, with generative AI applications such as ChatGPT being among the most used.

Table 1. Students’ profile

Demographic information		Frequency	Percentage (%)
Gender	Male	162	34.7
	Female	305	65.3
Age	18–19 years old	194	41.5
	20–21 years old	222	47.5
	22–23 years old	30	6.4
	24 years old and above	21	4.5
Level of study	Foundation	35	7.5
	Certificate	1	0.2
	Diploma	364	77.9
	Bachelor’s Degree	65	13.9
	Postgraduate (Master’s/PhD)	2	0.4
Institution	UCYP	125	26.8
	UniPSAS	84	18
	UiTM Pahang	91	19.5
	IIUM	31	6.6
	UMPSA	1	0.2
	KPTM Kuantan	117	25.1

(Continued on next page)

Table 1. (Continued)

Demographic information		Frequency	Percentage (%)
Race	KYP	6	1.3
	KKTM Kuantan	12	2.6
	Malay	454	97.2
	Chinese	3	0.6
	Other	10	2.1
Religion	Buddhism	4	0.9
	Christianity	5	1.1
	Islam	457	97.9
	Other	1	0.2
Type of AI used	Canva SlideALio	65	13.9
	ChatGPT	210	45.0
	Google Classroom	142	30.4
	Other	50	10.7

Descriptive Analysis of Ethical AI Use Among Students

Descriptive analysis was conducted to examine students' perceptions, emotions, and behaviours related to AI use. The analysis focused on four constructs, including perceived risk, verification intention, negative emotion, and responsible use of AI. The findings are summarised in Tables 2–5.

Perceived risk

Table 2 summarises the descriptive statistics for the Perceived Risk construct. Overall, the results indicate that students recognise the potential risks associated with AI-generated information, including misinformation, bias, and data privacy.

Table 2. Summary of perceived risk

Perceived risk	Frequencies					Total
	Strongly disagree	Disagree	Neutral	Agree	Strongly agree	
Frequent use of AI applications diminishes my ability to think critically and solve problems independently.	12 (2.6%)	47 (10.1%)	149 (31.9%)	175 (37.5%)	84 (18.0%)	467 (100%)
Frequent use of AI applications threatens the privacy and security of my personal data.	13 (2.8%)	53 (11.3%)	162 (34.7%)	173 (37.0%)	66 (14.1%)	467 (100%)

Verification intention

Table 3 presents the descriptive statistics for the Verification Intention construct. The results suggest that students generally demonstrate a strong intention to verify AI-generated information before incorporating it into their academic work.

Table 3. Summary of the verification intention

Verification intention	Frequencies					Total
	Strongly disagree		Neutral	Agree	Strongly agree	
I am willing to verify the information obtained from AI applications through additional sources before considering it as completely accurate or true.	7 (1.5%)	13 (2.8%)	127 (27.2%)	195 (41.8%)	125 (26.8%)	467 (100%)
I am willing to corroborate the information provided by AI applications through comparison with academic sources and experts in the relevant field.	4 (0.9%)	16 (3.4%)	133 (28.5%)	203 (43.5%)	111 (23.8%)	467 (100%)
It is not necessary to verify the veracity of the information provided by AI applications because it always provides valid and reliable information.	81 (17.3%)	147 (31.5%)	134 (28.7%)	82 (17.6%)	23 (4.9%)	467 (100%)

Negative emotion

Table 4 reports the descriptive statistics for the Negative Emotion construct. The findings indicate that some students have concerns about excessive reliance on AI technologies, particularly regarding their potential impact on independent thinking and academic integrity.

Table 4. Summary of the negative emotion

Negative emotion	Frequencies					Total
	Strongly disagree	Disagree	Neutral	Agree	Strongly agree	
I dislike the idea that technology such as AI applications replaces certain human skills such as inferring, information seeking, analysing, writing, etc.	16 (3.4%)	58 (12.4%)	212 (45.4%)	134 (28.7%)	47 (10.1%)	467 (100%)
I am concerned that frequent use of AI applications will limit my ability to think and solve problems independently.	11 (2.4%)	29 (6.2%)	138 (29.6%)	184 (39.4%)	105 (22.5%)	467 (100%)
I am concerned that excessive use of AI applications will diminish my interest in researching and reading diverse sources of information.	15 (3.2%)	43 (9.2%)	133 (28.5%)	169 (36.2%)	107 (22.9%)	467 (100%)

Responsible use of AI

Table 5 presents the descriptive statistics for Responsible Use of AI. Overall, students report relatively high levels of responsible AI usage, including verifying AI-generated outputs and using AI tools primarily as supportive learning resources.

Table 5. Summary of the responsible use of AI

Responsible use of AI	Frequencies					Total
	Strongly disagree	Disagree	Neutral	Agree	Strongly agree	
When using AI applications, I carefully check the responses to ensure that they are correct, complete, and free of bias.	3 (0.6%)	7 (1.5%)	99 (21.2%)	234 (50.1%)	124 (26.6%)	467 (100%)

(Continued on next page)

Table 5. (Continued)

Responsible use of AI	Frequencies					
	Strongly disagree	Disagree	Neutral	Agree	Strongly agree	Total
I use AI applications responsibly and ethically, ensuring that the answers obtained are a tool to support my learning and not a replacement for my own intellectual effort.	6 (1.3%)	5 (1.1%)	120 (25.7%)	222 (47.5%)	114 (24.4%)	467 (100%)
I use AI applications responsibly and ethically, avoiding the generation of misleading, false and/or biased content.	5 (1.1%)	4 (0.9%)	105 (22.5%)	233 (49.9%)	120 (25.7%)	467 (100%)
I strive to understand the limitations of AI applications and its potential to generate incorrect or biased responses, which motivates me to use it with caution and discernment.	4 (0.9%)	3 (0.6%)	110 (23.6%)	234 (50.1%)	116 (24.8%)	467 (100%)

Measurement Model Evaluation

The reliability of the constructs was assessed using Cronbach's alpha to evaluate the internal consistency of the measurement items. As presented in Table 6, the constructs Verification Intention ($\alpha = 0.828$), Negative Emotion ($\alpha = 0.831$), and Responsible Use of AI ($\alpha = 0.920$) demonstrate good internal consistency, exceeding the commonly accepted threshold of 0.70 (Hair et al., 2017). These results indicate that the items within each construct are consistently measuring the intended concept.

Table 6. Reliability test (Cronbach's alpha)

Constructs	Measurement items	Cronbach's alpha	Number of items
Perceived Risk	E12, E13	0.546	3(2)
Verification Intention	K32, K33	0.828	3(2)
Negative Emotion	M38, M39, M40	0.831	3(3)
Responsible Use of AI	N41, N42, N43, N44, N45	0.920	5(5)

However, the Perceived Risk construct recorded a Cronbach's alpha value of 0.546, which falls below the recommended threshold. Generally, a Cronbach's alpha below 0.70 may raise concerns about internal consistency reliability. However, Cronbach's alpha should

be considered alongside other reliability and validity measures, especially in PLS-SEM analysis, where composite reliability and convergent validity offer further evidence of measurement quality.

As reported in Tables 7 and 8, the Perceived Risk construct demonstrated acceptable levels of composite reliability (CR = 0.788), exceeding the recommended threshold of 0.70, and average variance extracted (AVE = 0.658), which is above the recommended minimum of 0.50 (Hair et al., 2019). These results indicate that the measurement items share sufficient common variance and adequately represent the underlying construct.

Despite the acceptable CR and AVE values, the relatively lower Cronbach's alpha suggests that the internal consistency of the Perceived Risk items may be moderate. This may be partly due to the limited number of retained indicators representing the construct. Therefore, although the construct was retained for further analysis because of its satisfactory validity indicators, the findings related to Perceived Risk should be interpreted with caution. Future studies might consider refining or expanding the measurement items to enhance the internal consistency of this construct.

Table 7. Convergent validity of measurement model

Constructs	Items	Loadings	CR	AVE
Perceived Risk	E12	0.639	0.788	0.658
	E13	0.953		
Verification Intention	K32	0.923	0.921	0.853
	K33	0.924		
Negative Emotion	M38	0.664	0.894	0.743
	M39	0.951		
	M40	0.940		
Responsible Use of AI	N41	0.826	0.940	0.76
	N42	0.810		
	N43	0.903		
	N44	0.931		
	N45	0.885		

Furthermore, the Fornell–Larcker criterion presented in Table 8 confirms that the square root of the AVE for Perceived Risk is greater than its correlations with the other constructs, indicating satisfactory discriminant validity. This result suggests that the Perceived Risk construct is empirically distinct from the other variables included in the model. Therefore, although the Cronbach's alpha value is relatively modest, the construct still satisfies the criteria for construct validity based on composite reliability (CR), average variance extracted (AVE), and discriminant validity. As suggested by scholars such as Hair et al. (2017) and Henseler et al. (2016), when other indicators of construct validity demonstrate acceptable

values, a slightly lower Cronbach’s alpha may still be considered acceptable, particularly in exploratory or developing research contexts. Accordingly, the Perceived Risk construct was retained for further analysis in this study.

Table 8. Discriminant validity of the measurement model

Constructs	Perceived risk	Verification intention	Negative emotion	Responsible use of AI
Perceived Risk				
Verification Intention	0.479			
Negative Emotion	0.549	0.243		
Responsible Use of AI	0.548	0.684	0.312	

The measurement model evaluation diagram (Figure 1) visually summarises these results, confirming the adequacy of the measurement properties.

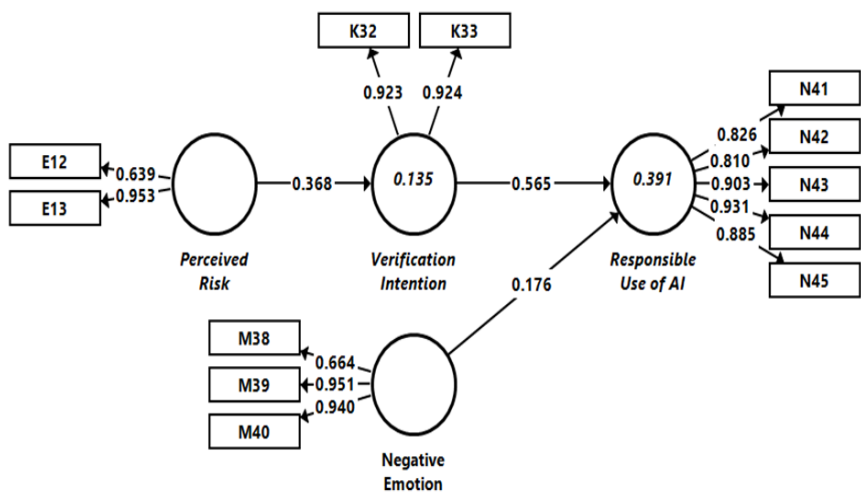


Figure 1. Measurement model evaluation diagram

Structural Model Evaluation

The structural model was evaluated by examining the path coefficients, t-values, and p-values for each hypothesised relationship. The results are summarised in Table 9 and visually depicted in Figure 2, which presents the full path analysis of the structural model.

Table 9. Path coefficient and hypothesis testing

Hypothesis	Relationship	Coefficient	<i>t</i> -value	Decision
H1	Perceived Risk → Verification Intention	0.368	7.327 (0.000)	Supported
H2	Verification Intention → Responsible Use of AI	0.541	12.449 (0.000)	Supported
H3	Negative Emotion → Responsible Use of AI	0.200	4.747 (0.000)	Supported
H4	Verification Intention → Negative Emotion → Responsible Use of AI (Moderation)	-0.109	2.238 (0.025)	Supported

H1: Perceived Risk → Verification Intention

The path coefficient from Perceived Risk to Verification Intention was positive and significant $\beta = 0.368$, $t = 7.327$, $p < 0.001$. This finding supports H1, indicating that students who perceive higher risks related to AI applications are more likely to intend to verify information obtained from AI sources before usage. This suggests that risk perception is an important driver of cautious and evaluative behaviour among students when engaging with AI-generated content.

H2: Verification Intention → Responsible Use of AI

The relationship between Verification Intention and Responsible Use of AI was also positive and statistically significant $\beta = 0.541$, $t = 12.449$, $p < 0.001$, thereby supporting H2. This result indicates that students who actively intend to verify AI outputs are more likely to use AI applications responsibly, such as by validating information, properly attributing content, and ensuring AI complements rather than replaces their intellectual efforts.

H3: Negative Emotion → Responsible Use of AI

A direct effect from Negative Emotion to Responsible Use of AI was tested in H3. The path was positive and significant $\beta = 0.200$, $t = 4.747$, $p < 0.001$, supporting H3. This result suggests that students who experience stronger negative emotions, such as concern, discomfort, or anxiety about AI misuse, are more inclined to engage in responsible AI usage. Negative emotional responses, therefore, act as a motivator that directly encourages ethical and cautious behaviour toward AI applications.

H4: Moderating Effect of Negative Emotion

The moderating effect of Negative Emotion on the relationship between Verification Intention and Responsible Use of AI was tested in H4. The moderation effect was statistically significant $\beta = -0.109$, $t = 2.238$, $p = 0.025$, thus supporting H4. Interestingly,

the negative coefficient suggests that while Negative Emotion strengthens the intention to act responsibly, at higher levels of Negative Emotion, the direct influence of Verification Intention on Responsible Use may weaken slightly. This could imply that excessive emotional concern might overwhelm rational verification processes or reflect an overcautious stance that alters normal verification behaviour. The interaction effect of Negative Emotion as a moderator is illustrated through the interaction pathway shown in Figure 2.

Overall, the structural model achieved strong empirical support, with all hypothesised relationships being statistically significant. The path diagram in Figure 2 provides a comprehensive overview of these results, visually demonstrating the strength and direction of each relationship in the proposed framework.

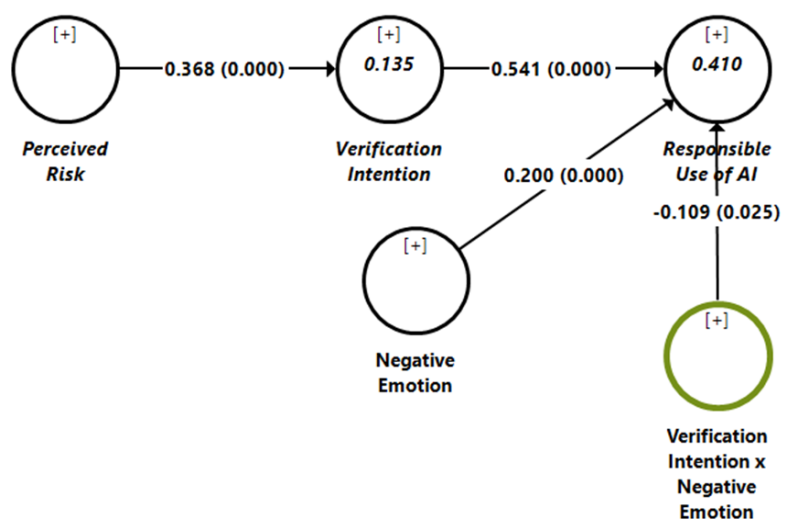


Figure 2. Result of path analysis

DISCUSSION

This study examined the factors influencing responsible AI use among university students by analysing the roles of perceived risk, verification intention, and negative emotion. The findings provide important insights into how students regulate their behaviour when interacting with AI-generated information in academic contexts.

The results indicate that perceived risk positively influences verification intention. This suggests that students who recognise the potential limitations of AI-generated information are more likely to verify AI outputs before using them in academic work. This finding highlights the role of risk awareness in encouraging responsible engagement with AI technologies. However, it is important to consider whether such verification behaviour

reflects genuine critical engagement or merely procedural checking. Prior research indicates that while AI tools can enhance efficiency, they may also reduce students' critical engagement and encourage reliance on generated outputs rather than independent thinking (Akgun & Greenhow, 2022). Students may therefore verify AI-generated information without fully evaluating its credibility or contextual relevance, raising concerns about the depth of their critical thinking. This suggests that verification may serve as a surface-level strategy to ensure compliance with academic expectations rather than promote meaningful learning.

The findings also show that verification intention significantly contributes to responsible use of AI. Students who intend to verify AI-generated information are more likely to engage in ethical practices such as cross-checking sources and ensuring proper use of AI outputs. This aligns with previous studies highlighting the importance of awareness of academic integrity in AI-assisted learning environments (Cotton et al., 2023). However, while verification behaviour appears to support responsible AI usage, it does not necessarily guarantee deeper learning or independent thinking. Students may rely on verification as a compensatory mechanism while still depending heavily on AI-generated content. This indicates that responsible behaviour may coexist with underlying dependency on AI tools, suggesting a more complex relationship between ethical practices and actual learning outcomes.

In addition, negative emotion was found to influence the use of responsible AI. Students who express concern or discomfort regarding AI usage appear more likely to adopt cautious and responsible behaviours. This supports prior research indicating that ethical concerns and uncertainty surrounding AI can shape user behaviour in educational contexts (Dwivedi et al., 2023). However, excessive concern may also lead to over-cautious or inconsistent use of AI technologies, potentially limiting the effective integration of AI as a learning tool. Emotional responses may therefore not always lead to optimal learning behaviours but instead introduce variability in how students engage with AI tools. This highlights the importance of balancing ethical awareness with constructive engagement in AI-supported learning.

From a theoretical perspective, the findings extend the TAM by applying it to the context of post-adoption ethical behaviour. While TAM traditionally explains technology adoption based on perceived usefulness and ease of use (Davis, 1989), this study demonstrates that behavioural intention remains relevant in explaining how users regulate their actions after adoption. However, the findings also suggest that TAM may not fully capture the complexity of ethical AI behaviour. Ethical decision-making involves not only behavioural intention but also cognitive evaluation of risks and emotional responses toward technology use. This indicates that extending TAM to post-adoption contexts requires integrating broader psychological and ethical dimensions. Furthermore, alternative perspectives, such as self-regulation or ethical decision-making frameworks, may offer additional insights into how students manage AI use in academic environments.

From an educational perspective, these findings suggest that universities should strengthen AI literacy and responsible AI practices among students. Emphasis should be placed on developing critical evaluation skills, encouraging independent thinking, and fostering awareness of the ethical implications of AI use. As AI technologies become increasingly embedded in academic environments, promoting responsible use is essential to ensure that these tools enhance rather than undermine learning outcomes.

Finally, although this study focuses on university students in Pahang, Malaysia, the findings contribute to broader discussions on ethical AI engagement in higher education. As AI technologies become increasingly integrated into educational environments worldwide, understanding how students interact with these tools responsibly is essential for maintaining academic integrity and promoting meaningful learning. In addition, this study contributes to the literature by shifting the focus from AI adoption to post-adoption ethical behaviour. By integrating cognitive, behavioural, and emotional factors, the study provides a more comprehensive understanding of how students regulate their use of AI technologies in academic settings.

CONCLUSION

In summary, this study highlights the importance of understanding the ethical use of AI from a post-adoption behavioural perspective. Moving beyond traditional technology adoption models, the findings demonstrate that responsible AI engagement among students is shaped by the interplay of cognitive evaluation (perceived risk), behavioural intention (verification intention), and affective response (negative emotion). This integrated perspective extends the application of the Technology Acceptance Model (TAM) by emphasising how users regulate their behaviour after adopting AI technologies, particularly in academic contexts where ethical considerations are critical.

By fostering this dual capacity, educators can develop students who are both technologically competent and ethically responsible. These individuals will be more likely to engage with AI in ways that reflect thoughtful judgement and ethical consideration for example, when the use of AI might constitute plagiarism, perpetuate biases, or compromise data privacy. Ultimately, this integrated approach helps prepare students to become not just competent users of technology but also conscientious digital citizens who actively contribute to a more ethical and equitable digital society.

Practically, educators could integrate structured AI literacy and ethics modules into coursework to promote informed and responsible use. Universities may consider establishing clear institutional guidelines, assessment policies, and verification procedures to support academic integrity when AI tools are used. Policymakers should work toward developing national standards for ethical AI use in education, with attention to data protection, transparency, and equitable access. Collectively, these coordinated measures would strengthen ethical AI practices within Malaysian higher education.

Future research could expand on the findings of this study by exploring the ethical use of AI technology applications across a broader demographic, including students from different states or regions in Malaysia for comparative analysis. Additionally, longitudinal studies could be conducted to examine how students' ethical awareness and behaviours evolve over time as AI technologies continue to advance and become more integrated into academic and personal settings. Future research could also incorporate qualitative or mixed methods approaches to gain deeper insights into students' motivations, perceptions, and ethical dilemmas when using AI tools. Furthermore, studies focusing on the role of institutional policies, digital ethics education, and the influence of cultural or religious values on ethical decision-making in AI use would contribute to a more comprehensive understanding of the issue.

LIMITATIONS

This study is not without limitations. First, the sample consisted predominantly of Malay (97.2%) and Muslim (97.9%) students enrolled in higher education institutions in Pahang. While this composition reflects the demographic characteristics of the region, it limits the generalisability of the findings to more diverse populations and educational contexts. Future research should consider broader and more heterogeneous samples to enhance external validity. Second, although this study adopts a post-adoption behavioural perspective grounded in the TAM, it focuses primarily on ethics-related constructs rather than incorporating the full set of traditional TAM variables. As such, the study does not capture the complete technology adoption process. Future studies may integrate both adoption and post-adoption constructs within a unified framework to provide a more comprehensive understanding of AI usage behaviour. Finally, the Perceived Risk construct demonstrated a relatively low Cronbach's alpha value, indicating moderate internal consistency. Although the construct was retained based on acceptable composite reliability and convergent validity, this limitation suggests that the measurement items may not fully capture the multidimensional nature of perceived risk. Future research is encouraged to refine or expand the scale to improve its psychometric robustness.

ACKNOWLEDGEMENTS

The study was funded by UCYP University under the UCYP University Internal Research Grant 2024 with reference code UCYP/TNC PI: 01/2024 (26), entitled "Analisa Etika Penggunaan Aplikasi Teknologi AI dalam Kalangan Mahasiswa di Negeri Pahang".

REFERENCES

- Akgun, S., & Greenhow, C. (2022). Artificial intelligence in education: Addressing ethical challenges in K-12 settings. *AI and Ethics*, 2(3), 431–440. <https://doi.org/10.1007/s43681-021-00096-7>

- Bond, M., Khosravi, H., De Laat, M., Bergdahl, N., Negrea, V., Oxley, E., Pham, P., Chong, S. W., & Siemens, G. (2024). A meta systematic review of artificial intelligence in higher education: A call for increased ethics, collaboration, and rigour. *International Journal of Educational Technology in Higher Education*, 21(1), 4. <https://doi.org/10.1186/s41239-023-00436-z>
- Chen, J. J., & Lin, J. C. (2024). Artificial intelligence as a double-edged sword: Wielding the POWER principles to maximize its positive effects and minimize its negative effects. *Contemporary Issues in Early Childhood*, 25(1), 146–153. <https://doi.org/10.1177/14639491231169813>
- 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*, 60(2), 178–189. <https://doi.org/10.1080/14703297.2023.2190148>
- Creswell, J. W., & Creswell, J. D. (2018). *Research design: Qualitative, quantitative, and mixed methods approaches* (5th ed.). SAGE Publications.
- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319–340. <https://doi.org/10.2307/249008>
- Dong, X., Wang, Z., & Han, S. (2025). Mitigating learning burnout caused by generative artificial intelligence misuse in higher education: A case study in programming language teaching. *Informatics*, 12(1), 51. <https://doi.org/10.3390/informatics12020051>
- Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., ... & Wright, R. (2023). Opinion paper: “So what if ChatGPT wrote it?” Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *International Journal of Information Management*, 71, 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- Elmessiry, A., Elmessiry, M., & Elmessiry, K. (2023). Unethical use of artificial intelligence in education [Paper presentation]. *Proceedings of the 15th International Conference on Education and New Learning Technologies*, Palma, Spain, 3–5 July, pp. 6703–6707. <https://doi.org/10.21125/edulearn.2023.1764>
- Fadlelmula, F. K., & Qadhi, S. M. (2024). A systematic review of research on artificial intelligence in higher education: Practice, gaps, and future directions in the GCC. *Journal of University Teaching and Learning Practice*, 21(6), 146–173. <https://doi.org/10.53761/pswgbw82>
- Fahd, K., Venkatraman, S., Miah, S. J., & Ahmed, K. (2022). Application of machine learning in higher education to assess student academic performance, at-risk, and attrition: A meta-analysis of literature. *Education and Information Technologies*, 27(3), 3743–3775. <https://doi.org/10.1007/s10639-021-10741-7>
- Floridi, L. (2018). Soft ethics, the governance of the digital and the General Data Protection Regulation. *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 376(2133), 20180081. <https://doi.org/10.1098/rsta.2018.0081>

- Floridi, L., & Cows, J. (2022). A unified framework of five principles for AI in society. In S. Carta (Ed.), *Machine learning and the city: Applications in architecture and urban design* (pp. 535–545). Wiley. <https://doi.org/10.1002/9781119815075.ch45>
- Garrote Jurado, R., Pettersson, T., & Zwierewicz, M. (2023). Students' attitudes to the use of artificial intelligence. *Proceedings of 16th International Conference of Education, Research and Innovation*, Seville, Spain, 13–15 November, pp. 514–519. <https://doi.org/10.21125/iceri.2023.0191>
- Hair, J. F., Hult, G. T. M., Ringle, C. M., & Sarstedt, M. (2017). *A primer on Partial Least Squares Structural Equation Modeling (PLS-SEM)* (2nd Ed.). Sage Publications Inc.
- Hair, J. F., Risher, J. J., Sarstedt, M., & Ringle, C. M. (2019). When to use and how to report the results of PLS-SEM. *European Business Review*, 31(1), 2–24. <https://doi.org/10.1108/EBR-11-2018-0203>
- Hanh, N. V., & Duyen, N. T. (2025). AI-assisted academic cheating: A conceptual model based on postgraduate student voices. *Frontiers in Computer Science*, 7, 1682190. <https://doi.org/10.3389/fcomp.2025.1682190>
- Imran, M., Almusharraf, N., Abdellatif, M. S., & Abbasova, M. Y. (2024). *Artificial intelligence in higher education: Enhancing learning systems and transforming educational paradigms*. Khazar University Repository.
- Kasneci, E., Sessler, K., Krosse, H., & Bannert, M. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Instruction*, 86, 101791. <https://doi.org/10.1016/j.lindif.2023.102274>
- King, M. R., & ChatGPT. (2023). A conversation on artificial intelligence, chatbots, and plagiarism in higher education. *Cellular and Molecular Bioengineering*, 16(1), 1–2. <https://doi.org/10.1007/s12195-022-00754-8>
- Kuka, L., Hörmann, C., & Sabitzer, B. (2022). Teaching and learning with AI in Higher Education: A scoping review. In M. E. Auer, A. Pester, & D. May (Eds.), *Learning with technologies and technologies in learning* (pp. 551–571). Cham: Springer. https://doi.org/10.1007/978-3-031-04286-7_26
- Li, Y., Sha, L., Yan, L., Lin, J., Raković, M., Galbraith, K., Lyons, K., Gašević, D., & Chen, G. (2023). Can large language models write reflectively. *Computers and Education: Artificial Intelligence*, 4, 100140. <https://doi.org/10.1016/j.caeai.2023.100140>
- Long, D. Y., Wang, S., Md Rashid, S., & Lu, X. T. (2025). Artificial intelligence in higher education: A systematic review of its impact on student engagement and the mediating role of teaching methods. *Frontiers in Education*, 10, 1648661. <https://doi.org/10.3389/feduc.2025.1648661>
- Lund, B. D., Lee, T. H., Mannuru, N. R., & Arutla, N. (2025). AI and academic integrity: Exploring student perceptions and implications for higher education. *Journal of Academic Ethics*, 23(3), 1545–1565. <https://doi.org/10.1007/s10805-025-09613-3>
- Marengo, A., Pagano, A., Pange, J., & Soomro, K. A. (2024). The educational value of artificial intelligence in higher education: A 10-year systematic literature review. *Interactive Technology and Smart Education*, 21(4), 625–644. <https://doi.org/10.1108/ITSE-11-2023-0218>

- Martiniello, N., Asuncion, J., Fichten, C., Jorgensen, M., Havel, A., Harvison, M., Legault, A., Lussier, A., & Vo, C. (2021). Artificial intelligence for students in postsecondary education: A world of opportunity. *AI Matters*, 6(3), 17–29. <https://doi.org/10.1145/3446243.3446250>
- Maulana, A., Idroes, G. M., Kemala, P., Maulydia, N. B., Sasmita, N. R., Tallei, T. E., Sofyan, H., & Rusyana, A. (2023). Leveraging artificial intelligence to predict student performance: A comparative machine learning approach. *Journal of Educational Management and Learning*, 1(2), 64–70. <https://doi.org/10.60084/jeml.v1i2.132>
- Mustapha, R., Mustapha, S. N. A. M., & Mustapha, F. W. A. (2024). Students' misuse of ChatGPT in higher education: An application of the Fraud Triangle Theory. *Journal of Contemporary Social Science and Education Studies*, 4(1), 87–97. <https://doi.org/10.5281/zenodo.10912353>
- Pallant, J. (2020). *SPSS survival manual: A step by step guide to data analysis using IBM SPSS* (7th ed.). Open University Press.
- Ramirez, E. A. B., & Esparrell, J. A. F. (2024). Artificial intelligence (AI) in education: Unlocking the perfect synergy for learning. *Educational Process: International Journal*, 13(1), 35–51. <https://doi.org/10.22521/edupij.2024.131.3>
- Selvi, B. (2024). The role of generative artificial intelligence tools in enhancing academic writing. *Social Paradigm*, 8(1), 22–56.
- Susnjak, T., & McIntosh, T. R. (2024). ChatGPT: The end of online exam integrity? *Education Sciences*, 14(6), 656. <https://doi.org/10.3390/educsci14060656>
- Tindle, R., Pozzebon, K., Willis, R., & Moustafa, A. A. (2023). Academic misconduct and generative artificial intelligence: University students' intentions, usage, and perceptions. *PsyArXiv Preprints*, 13. <https://doi.org/10.31234/osf.io/hwkgu>
- Zawacki-Richter, O., Marín, V. I., Bond, M., & Gouverneur, F. (2019). Systematic review of research on artificial intelligence applications in higher education: Where are the educators? *International Journal of Educational Technology in Higher Education*, 16(1), 1–27. <https://doi.org/10.1186/s41239-019-0171-0>
- Zhao, L., Rahman, M. H., Yeoh, W., Wang, S., & Ooi, K. B. (2025). Examining factors influencing university students' adoption of generative artificial intelligence: a cross-country study. *Studies in Higher Education*, 50(12), 2646–2668. <https://doi.org/10.1080/03075079.2024.2427786>