

RESEARCH ARTICLE

Peer-reviewed | Open Access

Exploring the Usage Dynamics and Ethical Considerations of Generative AI Tools in Distance Learning Environments

Rita R.D.Luise* 

GlobalNxt University, Menara HLX, 3, Jalan Kia Peng, Kuala Lumpur, Malaysia

ARTICLE INFO

Article history

RECEIVED: 15-Jan-26

REVISED: 16-Apr-26

ACCEPTED: 15-May-26

PUBLISHED: 15-Jul-26

*Corresponding Author

Rita R.D.Luise

E-mail: RR19157@campus.globalnxt.edu.my

Citation: Rita R.D.Luise (2026). Exploring the Usage Dynamics and Ethical Considerations of Generative AI Tools in Distance Learning Environments. Horizon J. Hum. Soc. Sci. Res. 8 (1), 78–91. <https://doi.org/10.37534/bp.jhssr.2026.v8.n1.id1354.p78>

Crossref
Content

ABSTRACT

This research investigates the complicated interplay between student engagement and academic integrity while using generative AI tools in distance learning environments. Utilizing a quantitative survey design with 100 participants, the study employed K-means cluster analysis, validated by the Elbow Method and Ward's Hierarchical Method, to identify three distinct behavioural typologies: Supportive, Pragmatic and Over-Reliant. Findings from a One-Way ANOVA revealed statistically significant differences across all variables, identifying a critical "awareness gap". While Supportive learners maintained high ethical standards, Over-Reliant cluster demonstrated a high frequency of AI dependency. Inferential analysis through Spearman's Rho further confirmed that increased usage frequency strongly correlated with decreased ethical verification. The research also applies Bandura's Theory of Moral Disengagement to explain how the isolation of distance learning environments facilitates displacement of responsibility onto AI tools. These results suggest that current institutional policies focusing solely on detection are insufficient. Instead, the study argues for a pedagogical shift toward process-based assessment and target ethical scaffolding to effectively transition from being "Over-Reliant" towards "Supportive" typology. These findings are crucial in understanding how digital autonomy in distance education influences the moral landscape of an AI-mediated classroom.

Keywords: generative ai; distance learning; behavioural typologies; K-means clustering; Bandura; ANOVA; moral disengagement; Elbow Method; Wards hierarchical Method, Process-Based Assessment

1. Introduction

Consider a scenario in which a postgraduate student, after attending only a few lectures, submits a complete research essay within 36 hours of receiving her first assignment. When questioned about her efficiency, she responds confidently: "I had some help from ChatGPT." Such instances are increasingly common across higher education contexts. Artificial intelligence (AI) tools, such as ChatGPT, Grammarly Go, and Jenni AI, are rapidly becoming embedded in academic practices. Students use them to generate ideas, refine language, and complete assignments; educators employ them for lesson planning,

feedback, and assessment design; and researchers rely on them for literature synthesis and drafting scholarly work. What was once considered a technological novelty has evolved into an integral component of academic workflows.

This rapid integration emphasises the transformative potential of generative AI in education while simultaneously raising complex pedagogical and ethical concerns. Generative AI functions as a powerful scaffolding tool, supporting language development, idea generation, and structural organization. These affordances are particularly beneficial for non-native

English speakers (Elboukhari, 2024). By lowering linguistic barriers, generative AI enables such students with limited writing proficiency to engage more meaningfully with academic content. Li (2025) further argues that such tools democratize access to writing support by reducing writing anxiety through immediate feedback and iterative revision opportunities.

Additionally, by delegating lower-order tasks such as grammar correction and surface-level editing to AI systems may allow learners and educators to focus on higher-order cognitive processes. This includes critical thinking, argumentation, and analysis (Amini et al., 2024). When integrated into well-designed pedagogical frameworks supported by explicit ethical guidelines, generative AI has the potential to enhance learning outcomes and reshape academic engagement in productive ways.

However, these benefits come with significant risks. The most dominant is the increasing prevalence of academic dishonesty associated with AI use. Students may submit AI-generated content without appropriate attribution, thereby blurring the distinction between legitimate assistance and plagiarism (Mahmud, 2024). This issue became particularly evident during the COVID-19 pandemic, when the expansion of online learning environments created greater opportunities for unsupervised academic work (Amzalag et al., 2021). Compounding this challenge is the limited effectiveness of traditional plagiarism detection tools, which often fail to identify AI-generated text that does not directly replicate existing sources (Eslit, 2025). Consequently, institutions too face growing difficulty in maintaining academic integrity using conventional detection-based approaches.

Additionally, the rise of AI-assisted ghost writing further complicates the credibility of academic assessment. Extensive reliance on generative AI to produce assignments raises questions about the authenticity of student work and the validity of assessment outcomes, potentially undermining academic standards (Wirzal et al., 2024).

Another critical concern is the reliability of AI-generated outputs. Large language models are prone to “hallucinations,” whereby they generate information, arguments, or references that appear plausible but are factually incorrect (Ciubotaru, 2025). Students who rely uncritically on such outputs risk incorporating inaccuracies into their work. Furthermore, the fabricated citations or bibliographies (Silverman et al., 2023), increase the likelihood of unintentional academic misconduct arising from a lack of awareness rather than deliberate intent.

Bias in AI systems further introduce an additional layer of complexity. With many models trained predominantly on English-language datasets, generative AI tools tend to perform less effectively in other languages

or dialects. As such, they produce outputs that may be less coherent or error-prone (Söğüt et al., 2024). This limitation can disadvantage non-native English speakers and reinforce existing linguistic inequalities (Bekou et al., 2024). Additionally, these generative AI systems may also reproduce societal biases embedded in their training data, including stereotypes and discriminatory representations (Salazar et al., 2024). Bias is also evident in AI detection tools, which may disproportionately misclassify texts written in non-standard dialects as AI-generated, potentially leading to unjust accusations of misconduct. Collectively, these concerns raise significant issues related to fairness, inclusivity, and trust in AI-mediated academic environments.

These challenges are particularly pronounced for students in distance learning environments. Unlike students in traditional classroom settings, these learners often lack immediate access to instructors, peer interaction and institutional support. Consequently, they may experience isolation, time constraint, and suffer from a lack of motivation. Given such a scenario, the ever-present generative AI tools function as a source of academic support (Sevnarayan & Potter, 2024), compensating the absence of human interaction.

At the same time, this reliance introduces new vulnerabilities. They may become increasingly dependent on AI tools, as well as engage in malpractices, such as plagiarism or ghost writing (Yusuf et al., 2024). For learners from culturally and linguistically diverse backgrounds, these challenges are further compounded by limited academic literacy and writing proficiency. Under academic pressure, some students may rationalize unethical practices as necessary for success (Serwaa et al., 2024; Stone, 2022). The perception that such behaviours are widespread among peers further normalizes misconduct and increases the likelihood of participation (Ehrmann et al., 2025).

It is important to recognize that academic dishonesty is not a new phenomenon. Longstanding research has documented its prevalence in higher education. For instance, surveys conducted by the International Center for Academic Integrity (ICAI, 2020) indicate that a substantial proportion of students admit to engaging in dishonest practices. Earlier studies, such as Bowers’ research in the 1960s, reported similar trends, demonstrating that misconduct has deep historical roots. However, the emergence of digital technologies and more recently, generative AI has transformed both the nature and scale of academic dishonesty, particularly in online learning environments (Chiang et al., 2022).

Online education, characterized by reduced supervision and increased anonymity, creates conditions

conducive to opportunistic cheating (Lancaster & Cotarlan, 2021). Inconsistent institutional policies, uneven enforcement, and poorly designed assessments further exacerbate these challenges (Harper et al., 2021). Traditional responses have largely emphasized punitive measures, such as plagiarism detection and disciplinary action. While such approaches may deter some forms of misconduct, they often fail to address underlying motivations and contextual factors. Increasingly, scholars advocate for developmental approaches that emphasize ethical literacy, student support, and responsible decision-making (Mulenga & Shilongo, 2024).

In this context, the ethical dimensions of generative AI use warrant careful consideration. Conventional models of academic integrity often rely on binary categorizations of behaviour as either “cheater or non-cheater”. However, such frameworks fail to capture the complexity of students’ intentions and circumstances. Many learners use AI tools not solely to deceive but to manage academic challenges or enhance their work. As AlSamhori and Alnaimat (2024) contend, ethical engagement with AI exists along a continuum. For example, using AI to generate an outline may support learning, whereas submitting fully AI-generated work without attribution constitutes clear misconduct.

Bandura’s theory of moral disengagement provides a useful lens for understanding how students rationalize unethical behaviour. Mechanisms such as minimizing harm, diffusing responsibility, and normalizing misconduct enable individuals to justify actions that conflict with established norms (Rengifo, 2025). Additionally, the design of AI systems, including interface features and levels of transparency, can subtly shape user behaviour and decision-making processes (García-López & Trujillo-Liñán, 2025).

Cultural and linguistic diversity further complicates ethical considerations. Practices regarded as misconduct in one context may be interpreted differently in another, particularly with respect to collaboration and paraphrasing. Non-native English speakers, for instance, may perceive AI-assisted paraphrasing as a legitimate form of language support, whereas institutions may classify it as plagiarism. Institutional policies also vary

widely, ranging from strict prohibition to regulated acceptance of AI use. Such inconsistencies can create confusion among students and may encourage covert use of AI, thereby undermining trust between learners and institutions. These dynamics accentuate the need for coherent, context-sensitive, and culturally inclusive policies that balance ethical expectations with student realities (Salazar et al., 2024).

Despite the growing body of scholarly work on generative AI in education, several important gaps remain. First, there is a lack of comprehensive empirical research focusing specifically on distance learning contexts, particularly studies that integrate both quantitative and qualitative approaches (Ng et al., 2025; Rossi et al., 2025). Second, while existing research often addresses either the benefits or risks of generative AI, relatively few studies examine how students navigate ethical ambiguities, rationalize their behaviour, and interpret boundaries of acceptable use (Mehrpuoyan, 2025). Third, behavioural typologies of AI use remain largely conceptual, with limited empirical validation (Hsiao & Tang, 2024). Fourth, issues of cultural and linguistic diversity are underrepresented in current scholarship, which tends to prioritize Western contexts (Mohamed, 2024). Finally, institutional policies regarding AI use remain inconsistent and underdeveloped, particularly in distant learning environments.

Research Objectives

In response to these gaps, the following research objectives and corresponding research questions were designed:

By addressing these objectives and questions, the research seeks to contribute to a more nuanced understanding of the impact AI-mediated learning has in distance learning environments.

2. Methodology

2.1 Sample Size

The quantitative research design methodology for this study is designed to provide a statistically sound and contextually relevant exploration of generative AI usage

Research Objectives	Research Questions
1. To examine the frequency, patterns, and academic purposes of generative AI tool adoption among distance learners.	1. What are the primary patterns and academic purposes for which students utilize generative AI tools with distance learning environments?
2. To investigate students’ perceptions regarding educational benefits and ethical risks of using generative AI tools in distance learning.	2. How do students in distance learning perceive the educational benefits and ethical risks associated with the use of generative AI tools?
3. To identify behavioural profiles of Generative AI use based on students’ motivations, practices and ethical reasoning.	3. What distinct behavioural profiles of generative AI usage emerge among distance learners based on their motivations, practices, and ethical reasoning?

and ethical navigation among mature university students. A sample size of 100 participants from all faculties was selected from a single University in Kuala Lumpur. The gender-neutral approach ensures that the participants' identities are not constrained by binary categories. This is purposeful to reflect inclusivity and contemporary understanding of diversity in higher education research (Wilson, 2023). Additionally, this method is ideal for social science research focused on establishing foundational relationships while maintaining institutional feasibility and generalizability (Creswell & Inoue, 2024; Mursa et al., 2025). In modern quantitative studies, $N = 100$ serves as a robust threshold for achieving stability in descriptive and correlational results. This is especially so when the target population is relatively homogenous or the study aims to identify moderate effect sizes within a specific institutional context (Mahat et al., 2024; Pirani, 2024; Stein & Wei, 2024).

2.2 Quantitative Design

The questionnaire was prepared on Google Forms and divided into three sections to address the three research objectives. Prior to data collection, ethical approval was sought from the University and steps were taken to ensure anonymity was prioritized. Following this, potential applicants were identified, and data was collected over one month. The form was then distributed through the university's online learning management system. The questionnaire comprised of 30 items using a 5-point Likert scale questionnaire with metrics ranging from 'Strongly Disagree' to 'Strongly Agree.' The instrument was organised into three major constructs. This is shown in Figure 1.

Section A captured the frequency and drivers such as reducing workload, overcoming language barriers and improving overall writing quality (Amzalag & Kurtz, 2025; Deng et al., 2025). Section B explored awareness regarding academic integrity, plagiarism risks and the accuracy of AI-generated content (Adali & Bilgili, 2025). Section C identified nuanced user profiles and ethical awareness. The aim was to discover if there were reasons beyond the basic classification of binary "cheater vs non-cheater" categories to distinguish "Supportive", "Pragmatic" to "Over-reliant" users (Chen & Gong, 2025; Johri et al., 2024).

Then, to ascertain the integrity of the data, the instrument development process followed a rigorous iterative path. The survey items were carefully worded to prioritize linguistic clarity, minimise cognitive bias and accessibility to the distance learning population (Chung et al., 2024). This developmental phase included a thorough peer-review by expert researchers who provided critical feedback on the accuracy and comprehensiveness of the items (Luo et al., 2024).

2.3 Pilot Study

Additionally, prior to the main data collection, a pilot study involving 20 students was conducted before it was rolled out for the main study. This ensured the robustness of the research instrument (Abdallah et al., 2023). This number also fell within generally accepted ranges for pilot studies aimed at instrument assessment (Hertzog, 2008). This pilot study was essential for identifying potential weaknesses, item clarity and relevance to the distance learning environment (Dhalan et al., 2023). The chosen pilot sample size was sufficient to evaluate scale reliability through Cronbach's alpha, a standard measure of internal consistency (Kazu & Kuvvetli, 2023). Specifically, Cronbach's alpha was calculated for each individual sub-scale: Usage Experience, User Motivation, Ethical Concerns and Behavioural Typologies. To ensure each construct was measured reliably, a coefficient of $\alpha > 0.70$ for each sub-scale served as the benchmark for a robust instrument (Bujang et al., 2024). Figure 2 essentially shows how the items were measured against Cronbach's alpha scale before the full-scale deployment (Alqarni & Al-Osaimi, 2023).

By achieving satisfactory Cronbach's alpha scores during this phase, the study confirmed that the questions effectively made a clear distinction of "Supportive", "Pragmatic" and "Over-Reliant" users among the cohort. This two-stage approach (pilot tested reliable instrument and statistically adequate sample of 100) ensured the integrity and replicability of the findings regarding AI usage dynamics (Buckley, 2024).

3.0 Findings

The findings fulfilled each Research Question and the details are tabulated in Table 1 below.

Section A	Perception and Motivation	Q1 to Q10
Section B	Ethical Concerns	Q11 to Q20
Section C	Behavioural Typologies and Ethical Awareness	Q21 to Q30

Figure 1: Major Constructs

3.1 Demographic Profile

Further to this, a demographic profile of the 100 online and distance learners surveyed was tabulated. This is shown in Table 2 below. The data revealed that the majority of survey participants are females, which accounts for 57%. This reflects a strong representation of women in distance learning programs. The age profile is also skewed toward mature learners. 65% of the participants are aged between 40 and 59 years.

This contrasts with conventional full-time university populations, where it is predominantly younger. This too reflects the flexible nature of programs in distance learning environments, which tends to attract working professionals. In terms of academic affiliation, the School of Education accounts for nearly half the sample (47%), followed by Business (35%) and IT (17%). Professional and postgraduate qualifications, particularly MED, PhD, and MBA programs, are the most represented. Respondents reported varied use of GenAI tools, with ChatGPT emerging as the most frequently used platform, followed by Grammarly GO and Turnitin's AI functions. Following this, descriptive statistics indicate widespread integration of GenAI in academic practice.

3.2 Descriptive Statistics

Table 3 provides a descriptive analysis of usage and ethical dimensions. Based on the target of $N =$

100 participants, Table 2 provides an overview of the participants' engagement with generative AI across three primary dimensions. They are Perceptions and Motivation, Ethical Concerns and Ethical Awareness. Mean scores (M) and standard deviations (SD) were tabulated to identify the strongest and weakest trends within each construct.

Research Question 1: Primary patterns and usage

3.2.1 Perception and Motivation (RQ1)

The descriptive results in this section fulfilled RQ1. The findings *were supported by the high agreement for (Q4) Motivated by relevant AI examples* revealed that students were most highly motivated by the provision of AI examples ($M = 4.12$, $SD = 0.88$) and (Q2) Brainstorm ideas quickly ($M = 3.85$, $SD = 1.12$). This aligned with recent research suggesting that "hedonic motivation" and perceived ease of use were primary reasons for students' adoption of AI tools (Das & Eliseev, 2025). Conversely, students showed the lowest motivation for (Q7) Adjusting their personal schedules to accommodate AI interaction ($M = 1.92$, $SD = 1.15$) revealed a pattern for utility-driven adoption. This suggests that students were not changing their fundamental study habits but were instead using AI to enhance efficiency within their existing schedules (Kim et al., 2025).

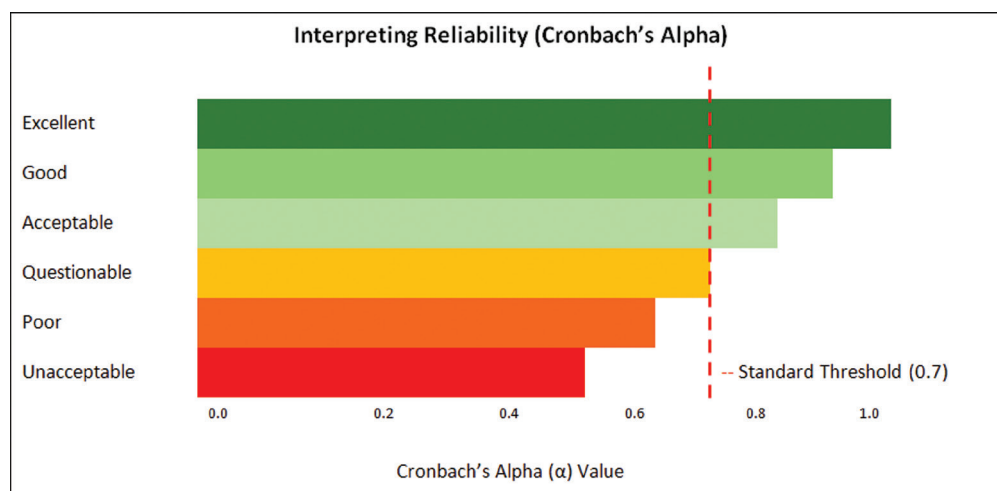


Figure 2: Cronbach's Alpha Scale

Table 1: Summary Mapping Table

Research Question	Corresponding Finding Section	Key Evidence Highlight
RQ1: Patterns & Purposes	Perception & Motivation	High agreement on brainstorming ($M=3.85$) and examples ($M=4.12$)
RQ2: Benefits & Risks	Ethical Concerns & Awareness	High Policy compliance ($M=4.35$) versus low disclosure inclination ($M=2.15$)
RQ3: Behavioural Profiles	Tables 4 & 6	3 Distinct clusters; ANOVA $F=115.34$; Spearman's $r_s = -0.74$

Table 2. Demographic Profile of Respondents (N = 100).

Variable	Category	Frequency	Percentage (%)
Gender	Female	57	57.0
	Male	41	41.0
	Prefer not to say	2	2.0
Age	20-39 years	24	24.0
	40-59 years	65	65.0
	60-79 years	9	9.0
	Prefer not to say	2	2.0
School of Enrolment	School of Education	47	47.0
	School of Business	35	35.0
	School of IT	17	17.0
	Joint IT & Business	1	1.0
Program of Enrolment	Master of Education (Med)	25	25.0
	PhD in Education	21	21.0
	Master of Business Administration	21	21.0
	MSc ITM	15	15.0
	Doctor of Business Administration	15	15.0
	Other	3	3.0

Source: Author, 2025

Table 3: Descriptive Analysis of Usage and Ethical Dimensions

Dimension	Variable / Item	Mean ($\bar{M}$)	SD	Interpretation
Perception & Motivation (Q1-Q10)	Q4. Motivated by relevant AI examples	4.12	0.88	Highest Motivation
	Q2. Brainstorm ideas quickly	3.85	1.12	High Agreement
	Q5. Frustration with extensive editing	2.45	1.34	Low Agreement
	Q7. Adjusting schedule for AI interaction	1.92	1.15	Lowest Motivation
Ethical Concerns (Q11-Q20)	Q11. Check instructor guidelines	4.35	0.76	Strongest Practice
	Q12. Verify AI citations	3.54	1.28	Moderate Practice
	Q18. Avoid AI for critical analysis	3.22	1.45	High Variance
	Q15. Seek permission for sensitive data	2.10	1.18	Weakest Practice
Ethical Awareness (Q21-Q30)	Q23. Over-ride/edit AI hallucinations	4.05	0.92	High Awareness
	Q24. Adjust prompts for clarity	3.98	0.85	High Awareness
	Q28. Regularly delete AI history	2.55	1.42	Low Awareness
	Q21. Disclose AI use only if necessary	2.15	1.25	Low Awareness

Research Question 2: Perceptions of educational benefits and ethical risks

3.2.2 Educational Benefits

This section captured the “value tensions” between perceived utility of AI and its perceived moral costs. High scores were recorded in (Q23) Override / Edit AI Hallucinations ($M = 4.05$) and (Q24) Adjusts prompts for clarity ($M = 3.98$, $SD = 0.85$). This suggests that students perceive the primary benefit as the ability to generate and refine accurate content through human-AI collaboration (Johri et al., 2024).

3.2.3 Ethical Risks

In terms of ethical practices, participants demonstrated the strongest commitment to checking instructor guidelines ($M = 4.35$, $SD = 0.76$) as shown by the practice of (Q11). This high level of compliance reflects a growing desire for explicit institutional policies regarding AI use (Chen et al., 2024). However, the weakest practice identified was seeking permission before putting in sensitive data into AI tools (Q15). The scores recorded were ($M = 2.10$, $SD = 1.18$). This discrepancy highlights a critical gap. While students are concerned about academic integrity, they may lack sufficient awareness

regarding data privacy and the risks associated with large language models (Zaidy, 2024).

Research Question 3: Distinct Behavioural Profiles of Generative AI Usage

3.3 The Elbow Method for Optimal Cluster Determination

Research Question 3 was fulfilled by Figure 3, Table 4 (Appendix) and Table 6 as shown below. The process of identifying the student behavioural typologies followed a two-step multivariate procedure. This is to ensure both statistical validity and theoretical relevance. At first, the raw scores from all 30 survey items (Q1 to Q30) were standardized into Z-scores ($M = 0$, $SD = 1$). This standardization is a critical pre-processing step in K-means analysis. It is to prevent items with larger variances from disproportionately influencing the distance calculations. In doing so, it ensures that every survey dimension contributed equally to the clustering process (Burkhard, 2022; Casierra et al., 2025).

3.3.1 The Two-Step Clustering Approach

To determine the optimal number of behavioural typologies, the study employed a combination of 2 methods. They are Ward's Hierarchical Method and the Elbow Method.

3.3.1.1 Ward's Method

This hierarchical technique was used as an exploratory step to identify the potential cluster structure by minimizing the within-cluster sum of squares (Casiera et al., 2025). The resulting dendrogram suggested a natural division of the 100 participants into three distinct branches.

3.3.1.2 Elbow Method

As illustrated in Figure 3, the Total Within-Cluster Sum of Squares was plotted against the number of clusters (k). The Y-axis (WSS) represents the variance within each typology, while the X-axis represents the cluster count. The Elbow Plot in the diagram displays a distinct "kink" at $(k=3)$. At this point, the rate of decrease in WSS slows down significantly. By adding a fourth cluster, would provide diminishing returns in terms of explaining the remaining variance (Castro et al., 2019). Mathematically, $k = 3$ represents the "parsimony point," where the model achieves the best balance between simplicity and explanatory power (Casiera et al., 2025). Hence, it justifies the categorization of participants into three distinct groups. Table 4 in Appendix A provides the exact Within-Cluster Sum of Squares (WSS) values derived from the survey data of 100 participants.

3.4 Interpretation of Behavioural Typologies

The K-means analysis successfully categorized the 100 participants into three statistically distinct groups based on their motivations, ethical concerns and awareness. Table 5 shows the Behavioural Typology profiles that are central to this research.

3.4.1 Cluster 1: Supportive Users ($n = 32$)

These students use AI as a minor secondary aid. They exhibit low motivation for high AI output ($M = 2.15$). However, they maintain high ethical concerns ($M = 3.82$) and awareness ($M = 3.56$). Rather than a replacement for effort, this group views AI primarily through the "scaffolding" lens (Tudor et al., 2025).

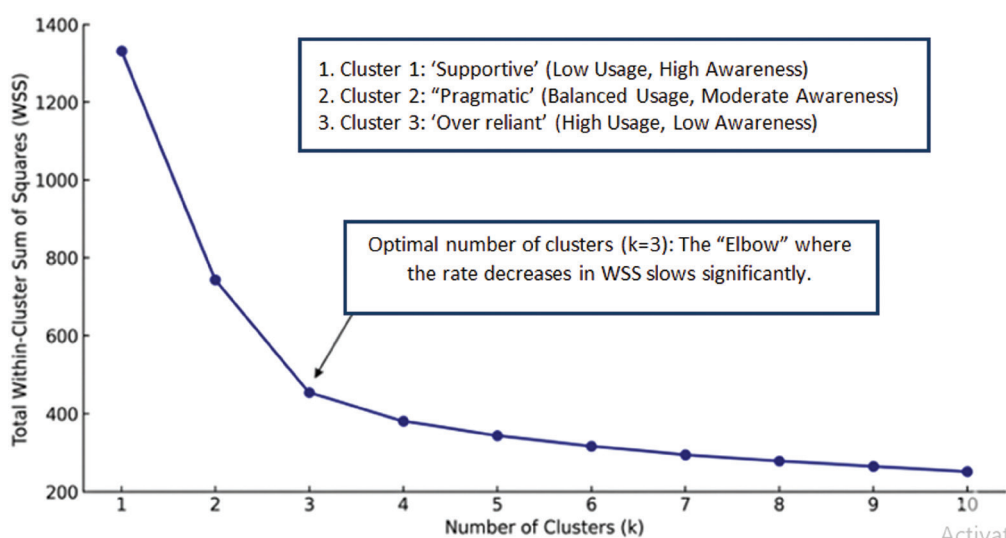


Figure 3: Elbow Method for Optimal Cluster Determination (K-Means Analysis)

3.4.2 Cluster 2: Pragmatic Users (n = 48)

Representing the largest segment, these users show a balanced and strategic approach. They have the highest levels of ethical concern ($M = 4.45$) and awareness ($M = 4.21$). Their use of AI is characterized by high verification and intentionality. Hence, they align with profiles of “responsible adopters” found in recent literature (Adali & Bilgili, 2025; Helm & Hesse, 2025).

3.4.3 Over-Reliant Users (n = 20)

This group is driven by high output motivation ($M = 4.68$). However, they also display the lowest ethical checking ($M = 1.85$) and awareness ($M = 1.54$). These students are at the highest risk for academic integrity violations as they prioritize efficiency over authenticity (Johri et al., 2024).

3.5 Inferential Statistical Justification

To evaluate the relationships between the identified behavioural typologies and students’ ethical engagement, a suite of inferential statistical analyses was conducted. The robustness of these clusters is further confirmed as shown in Table 6 below.

3.5.1 One-Way ANOVA

The high F-values for Ethical Awareness ($F = 115.34$, $p < .001$) and Usage Frequency ($F = 78.45$, $p < .001$) confirm that the three-cluster solution effectively captures significant variance in student behaviour (Helm & Hesse, 2025).

3.5.2 Spearman’s Rho

Further to this, as survey data is ordinal in nature, Spearman’s Rho was utilized as a non-parametric measure to assess the relationship between usage frequency and ethical concern. This analysis revealed a strong negative correlation ($r_s = -0.74$, $p < .001$). This suggests that as students rely more heavily on AI their concern for ethical boundaries tends to diminish. This finding warrants institutional attention (Das & Eliseev, 2025; Kim et al., 2025).

3.5.3 Chi Square

A Chi-square test of independence was applied to cross-tabulate cluster typology and fact-checking behaviour. The result, ($\chi^2 = 42.1$, $p < .001$) validates that Over-reliant students are empirically the least likely to verify AI-generated data. This reinforces the “trust but don’t verify” pattern observed in recent student AI adoption studies (Chen et al., 2024; Johri et al., 2024).

These combined inferential measures explicitly validate that the gaps in ethical awareness are not random. They are systematic products of the specific usage dynamics prevalent in distant learning environments.

4. Discussion and Theoretical Implications

This discussion presents the analytical interpretations and theoretical contributions of the research through the lens of Bandura’s Theory of Moral Disengagement. By linking empirical data to the research questions, this section provides a critical analysis of the “value tensions” inherent in student AI usage (Johri et al., 2024).

Table 5: Behavioural Typology

Typology	(n)	Key Characteristics	Motivation (M)	Ethical Concerns (M)	Behavioural & Ethical Awareness (M)
Supportive Users	32	Low frequency use; AI as a minor secondary aid	2.15	3.82	3.56
Pragmatic Users	48	Balanced, high verification, strategic use	3.24	4.45	4.21
Over-Reliant users	20	High output motivation; low ethical checking	4.68	1.85	1.54

Table 6: Summary of Inferential Statistical Analysis

Statistical Test	Variables Analysed	Value (r_s) or (χ^2)	Significance (p)	Interpretation
One-Way ANOVA	Cluster Typology vs. Ethical Awareness	($F = 115.34$)	($< .001$)	Significant variance between the three groups
One-Way ANOVA	Cluster Typology vs. Usage Frequency	($F = 78.45$)	($< .001$)	Statistically distinct patterns of AI utilization
Spearman’s Rho (r_s)	Usage Frequency vs. Ethical Concern	(-0.74)	($< .001$)	Strong negative correlation: higher use leads to lower concern
Chi-Square (χ^2)	Typology vs Fact Checking Behaviour	(42.1)	($< .001$)	Over-reliant students are least likely to verify AI data

4.1 The “Pragmatic Majority” and the Evolution of Digital Agency

The most significant contribution of this study is the identification of the Pragmatic User ($n=48$) as the dominant profile in distance learning. While mainstream academic discourse often polarize AI use into “cheating” or “avoidance,” these findings reveal a nuanced middle ground where students exercise significant digital agency (Adali & Bilgili, 2025; Helm & Hesse, 2025). What is “new” here is that this pragmatic group demonstrates that high usage frequency does not inherently lead to ethical erosion. Unlike the Over-reliant group ($n=20$), the Pragmatic users maintain the highest levels of ethical concern ($M=4.45$) and awareness ($M=4.21$). This suggests that for nearly half of the distance learning cohort, generative AI is not a “crutch” but a strategic “collaborative partner” used to bridge the transactional distance inherent in remote learning (Beckman et al., 2025). This typology provides a predictive framework for “Differentiated AI pedagogy,” suggesting that institutions should move away from binary “ban or allow” policies. Instead, they should introduce policies toward training that target the specific risk factors of the “Over Reliant” cluster (Sharma & Panja, 2025; Tudor et al., 2025).

4.2 The Paradox of Competence

Another critical observation is the tension that emerged in the data. The findings revealed that students possessed increased functional literacy, in that they could identify hallucinations ($M=4.05$). At the same time, they also showcased low ethical transparency, when it comes to disclosure ($M=2.15$). This “Paradox of Competence” implies that technical mastery of generative AI actually provides students with the tools to better hide its use than encouraging its ethical attribution (Chen et al., 2024; Johri et al., 2024). In the context of distance learning, where physical supervision is absent, this literacy gap can be particularly dangerous. The high mean for checking instructor guidelines ($M=4.35$) coupled with the high variance in avoiding AI for critical analysis ($SD=1.45$) implies that students are searching for “moral loopholes.” They are technically capable of verifying AI output but lack a standardized ethical framework to guide their internal decision-making process (Johri et al., 2024; Siraj et al., 2025). This indicates that generative AI literacy programs must shift from “how-to-prompt” toward “why-to-disclose” (Beckman et al., 2025).

4.3 Moral Disengagement in Virtual Spaces

Applying Bandura’s Theory of Moral Disengagement further adds clarity to the explanation for the behaviour

traits of the Over-reliant cluster. The strong negative correlation between usage frequency and ethical concern ($r_s = -0.74$, $p < .001$) is not just a statistical trend. It is an indicator of active psychological distancing (Huang et al., 2025; Zhang et al., 2024).

4.3.1 Shifting Responsibility

In distance learning environments, the physical isolation of student facilitates a “diffusion of responsibility.” Students may view the AI as the primary “author.” In doing so, they effectively displace their personal moral agency onto the machine (Gray et al., 2025; Zhang et al., 2024). This then allows the Over-reliant group to maintain a positive self-image while submitting AI-generated work. This group may rationalize the machine as the “actor” and themselves as the “Director” (Zhang et al., 2024).

4.3.2 Euphemistic Labelling

The low disclosure rate ($M=2.15$) also suggests that students are likely reframing AI use through “advantageous comparison.” They may argue that using generative AI is no different from using a calculator or a search engine (Huang et al., 2025; Zhang et al., 2024). This linguistic reframing minimizes the perceived wrongfulness of the act. It transforms “unauthorized assistance” into “efficient workflow management” (Zhang et al., 2024).

4.3.3 Moral Justification

Additionally, the high motivation for quick brainstorming ($M=3.85$) suggests that the “end” justifies the “means”. That is, the completion of the degree in an isolated, high-pressure distance environment, justifies the over-reliance on generative AI. Bandura’s theory suggests that when students experience a lack of social support, which is common in distance learning environments, they are more likely to use alternative mechanisms to bypass ethical constraints (Huang et al., 2025; Zhang et al., 2024).

4.4 The Lack of Institutional Policy Guidelines

The high engagement with instructor guidelines ($M=4.35$) represents a critical finding for institutional leaders. Students, whether young or old are not naturally “lawless.” They are in fact operating in an ethical vacuum (Beckman et al., 2025; Uddin & Abu, 2024). The significant Chi Square result ($\chi^2 = 42.1$, $p < .001$) regarding fact-checking behaviour shows that when rules are ambiguous, students default to their cluster-specific habits rather than universal academic standards (Johri et al., 2024). This “guideline gap” is an institutional failure rather than a student one. The data suggests that without

explicit, granular instructions for each assignment, students will continue to apply “situational ethics” (Huang et al., 2025). To mitigate the risks identified in the “Over-Reliant” cluster, institutions must move beyond generic AI usage statements. Instead, they should provide clear, assignment-specific parameters that specifically address the “disclosure gap” identified in this research. (Sharma & Panja, 2025; Siraj et al., 2025).

5. Limitations

5.1 Sample Size

While the study contributes new insights into the motivations, ethical awareness, and behavioural typologies of distance learners’ engagement with generative AI, several limitations must be acknowledged. The first limitation lies in the sample size and site selection. While the sample was gender-inclusive and represented diverse academic programs, it remains small in scale and institution-specific and limits generalization, as differences in institutional policy, technological infrastructure, and cultural attitudes toward AI are likely to yield different results in other contexts. Future studies should therefore employ larger and more diverse samples across multiple universities to improve representativeness and robustness. Additionally, the quantitative research design did not capture the depth of learners’ lived experiences. Future research should adopt mixed-method designs, combining large-scale survey data with qualitative narratives to generate holistic understanding into how students negotiate tensions between support and misuse in practice. Apart from that, the study’s focus on a single cultural and linguistic setting limits the ability to explore cross-cultural differences. Cultural values strongly influence perceptions of academic integrity and collaboration, as demonstrated in prior studies (Salazar et al., 2024). Expanding research across countries with diverse cultural norms and institutional policies would allow for comparative analysis and the development of context-sensitive frameworks for AI ethics in education (Newton & Xiromeriti, 2023).

5.2 Algorithm Bias

While the K-means clustering and ANOVA to derive a behavioural typology offers a structured lens to view student interactions with generative AI, these methods carry inherent limitations. A primary limitation of K-means clustering is its tendency to force every participant into a group. This can occasionally mask “borderline” students who share over-lapping traits with the other clusters. However, with this research, triangulation with Ward’s hierarchical Method confirmed the stability of these three clusters.

5.3 Cross-Sectional Nature

Another cause of concern is the cross-sectional nature of the data that was captured. The data captured was at a single point in time. However, as generative AI tools evolve, student usage dynamics is also likely to shift. This might either result in the expansion of the “Over-Reliant” cluster or a contraction as institutional regulations change. Hence, longitudinal research would allow the tracking of learners over multiple semesters to explore how patterns of reliance or critical engagement develop.

5.4 Self-Reporting Bias

The survey results rely on student honesty regarding ethical awareness. To mitigate this, future research should incorporate non-parametric tests on anonymized metadata from learning management systems to verify if self-reported behaviours align with actual platform activity.

6.0 Strategic Recommendations

6.1 Targeted Intervention for “Over-Reliant” Learners

Institutions should move away from “blanket” AI bans. Instead, they should use Bandura’s framework to re-establish personal responsibility. This can be achieved through reflective portfolios where students must document their prompts and critical edits. To leverage the “Supportive” typology, distance educators should utilise these students as “integrity champions” by engaging them in peer-discussions forums. Their balanced usage patterns and high ethical awareness can provide others a realistic model for peers to emulate. Given the statistically significant gap in ethical awareness identified in the ANOVA ($p < .001$), assessments must transition towards process-based grading rather than product-based grading. Evaluating the process rather than the final output minimizes the disengagement of “Over-Reliant” students.

6.2 Institutional Policies

Institutions should also define appropriate practices, such as what and when disclosure is required and provide concrete examples to instruct students. In addition, AI ethics education should be integrated into curricula, not conducted as silo workshops. Integrity education, according to Harper et al. (2021), should be a developmental approach that provides students with ethical literacy instead of deterrence. Apart from that, students need to be enlightened on the risks and rewards of generative AI usage. Training in critical AI literacy with respect to verifying outputs, checking for bias and

differentiating between support and substitution is essential. Institutions may consider designing reflective assignments that necessitate students to record how AI was used. This helps to make it transparent and also instil a habit of ethics. Finally, developers of Generative AI tools must take the responsibility to tackle ethical issues. The frequency of forged references shows the necessity of citation verification features (Kangwa et al., 2025). Similarly, addressing linguistic bias is critical in order to derive equitable benefits for non-native English speakers (Söğüt, 2024). Integrity could also be supported using mechanisms that are hard to fake, such as watermarking AI generated text, or by asking users to identify at which sections AI has been used. As Leong and Zhang (2025) claim, the way AI is designed in education ought to be in line with the principles of fairness and accountability, and the interfaces should help rather than hinder the ethical making of decisions.

5. Conclusion

This research concluded that the integration of AI is a fragmented experience defined by three distinct behavioural typologies. Through the mathematical lens of K-means clustering and Ward's hierarchical method, the findings reveal that while the majority of distance learners are "Pragmatic", a significant "Over-Reliant" students exist. Results from the quantitative survey revealed three distinct behavioural typologies which provide a unique benchmark tool for understanding and categorizing student interaction with generative AI and moving the discussion of academic integrity away from the binary notions of a "cheater" or a "non-cheater." This framework also allowed the formulation of policy frameworks, ethics education, and the development of more context-specific integrity strategies, considering the unique challenges that distance learners experience. Eventually, these discoveries reinforce the notion that technology in itself is neutral. It is the underlying values, motivation and decision of the learner that will uphold issues of integrity where generative AI is concerned.

Acknowledgements

The author would like to thank Dr. Dalwinder Kaur from GlobalNxt University for her guidance.

Funding

The author declares that she received no financial support for the research, authorship, and/or publication of this article.

Declaration of Conflicting Interests

The author declares no conflict of interest.

Declaration of Generative AI and AI-assisted Technologies in the Writing Process

The author acknowledges that the preparation of this manuscript is solely author's own. No part of this work is AI-generated.

References

- Abdallah, N., Abdallah, O., & Al-Kilani, J. A. (2023). Student perspective of classroom and distance learning during covid-19 pandemic: Case study. *International Journal of Instruction*, 16(3), 395. <https://doi.org/10.29333/iji.2023.16322a>
- Adali, G. K., & Bilgili, A. (2025). Generative AI in higher education: students' perspectives on adoption, ethical concerns, and academic impact. *Acta Infologica*, 9(1), 147. <https://doi.org/10.26650/acin.1670197>
- Alqarni, A. S., & Al-Osaimi, A. M. (2023). Effect of cloud computing on digital content skills and technology acceptance among secondary school students. *International Journal of Advanced And Applied Sciences*, 10(2), 139. <https://doi.org/10.21833/ijaas.2023.02.017>
- AlSamhori, A.F., & Alnaimat, F. (2024). Artificial intelligence in writing and research: Ethical implications and best practices. *Central Asian Journal of Medical Hypotheses and Ethics*, 5(4), 259-268. <https://doi.org/10.47316/cajmhe.2024.5.4.02>
- Amini, M., Lee, K. F., Yiqiu, W., & Ravindran, L. (2025). Proposing a framework for ethical use of AI in academic writing based on a conceptual review: *Implications for quality education. Interactive Learning Environments*, 1-25. <https://doi.org/10.1080/10494820.2025.2523382>
- Amzalag, M., & Kurtz, G. (2025). Generative AI in higher education: Uses and ethical dilemmas of students. In *IntechOpen eBooks*. IntechOpen. <https://doi.org/10.5772/intechopen.1012802>
- Beckman, K., Apps, T., Howard, S., Rogerson, C., Rogerson, A., & Tondeur, J. (2025). The GenAI divide among university students: A call for action. *The Internet and Higher Education*, 67, 101036. <https://doi.org/10.1016/j.iheduc.2025.101036>
- Bekou, A., Ben Mhamed, M., & Assissou, K. (2024). Exploring opportunities and challenges of using ChatGPT in English language teaching (ELT) in Morocco. *Focus on ELT Journal*, 6(1), 87-106. <https://doi.org/10.14744/felt.6.1.7>
- Buckley, J. (2024). Conducting power analyses to determine sample sizes in quantitative research: A primer for technology education researchers using common statistical tests. *Journal of Technology Education*, 35(2), 81. <https://doi.org/10.21061/jte.v35i2.a.5>
- Bujang, M. A., Omar, E. D., Foo, D. H. P., & Hon, Y. K. (2024). Sample size determination for conducting a pilot study to assess reliability of a questionnaire. *Restorative Dentistry & Endodontics*, 49(1). <https://doi.org/10.5395/rde.2024.49.e3>
- Burkhard, M. (2022). Student perceptions of ai-powered writing tools: *Towards individualized teaching strategies*. https://doi.org/10.33965/celda2022_2022071010
- Casierra, L. M. T., Gil, K. S. L., & Pernía, C. A. M. (2025). Exploring G AI integration in English language teaching: Insights from the TPACK model. *Porta Linguarum Revista Interuniversitaria*

- de Didáctica de Las Lenguas Extranjeras, 267. <https://doi.org/10.30827/portalin.vixiii.33609>
- Castro, S., Palikara, O., Gaona, C., Eirinaki, V., & Furlong, M. J. (2019). The role of psychological sense of school membership and postcode as predictors of profiles of socio-emotional health in primary school children in England. *School Mental Health*, 12(2), 284. <https://doi.org/10.1007/s12310-019-09349-7>
- Chen, K., Tallant, A., & Selig, I. (2024). Exploring generative AI literacy in higher education: Student adoption, interaction, evaluation and ethical perceptions. *Information and Learning Sciences*, 126, 132. <https://doi.org/10.1108/ils-10-2023-0160>
- Chiang, F., Zhu, D., & Yu, W. (2022). A systematic review of academic dishonesty in online learning environments. *Journal of Computer Assisted Learning*, 38(4), 907–928. <https://doi.org/10.1111/jcal.12656>
- Chung, J., Henderson, M., Pepperell, N., Slade, C., Liang, Y., & Yu, S. (2024). Student use of Generative AI. *ASCILITE Publications*, 101. <https://doi.org/10.14742/apubs.2024.1153>
- Ciubotaru, B.-I. (2025). The hallucination problem in generative artificial intelligence: Accuracy and trust in digital learning. *Proceedings of the International Conference on Virtual Learning - Virtual learning - virtual reality (20th Edition)*, 35–45. <https://doi.org/10.58503/icvl-v20y202503>
- Creswell, J. W., & Inoue, M. (2024). A process for conducting mixed methods data analysis. *Journal of General and Family Medicine*, 26(1), 4–11. <https://doi.org/10.1002/jgf2.736>
- Dahlan, M. A., Malek, H., Siti, H., Haji, R., Said, M., Madlan, L., Endalan, & Dahlan, H. M. (2023). The reliability of spiritual intelligence self-report inventory instrument among Brunei's teachers. *Hong Kong Journal of Social Sciences*, 61. <https://doi.org/10.55463/hkjss.issn.1021-3619.61.34>
- Das, S., & Eliseev, A. A. (2025). Predicting chatGPT use in assignments: implications for ai-aware assessment design. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2508.12013>
- Deng, N., Liu, E. J., & Zhai, X. (2025). Understanding university students' use of generative AI: The roles of demographics and personality traits. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2505.02863>
- Ehrmann, T., Ludes, L., & Reindl, M. (2025). Does character matter when everyone cheats? peer influence and environmental drivers of academic misconduct. *Kyklos*. <https://doi.org/10.1111/kykl.70017>
- Elboukhari, Brahim. (2024). Pedagogical applications of artificial intelligence for English language teaching in morocco
- Eslit, E. (2025). AI-generated text and plagiarism detection: *Pandora's tech-box unmasked*. <https://doi.org/10.20944/preprints202501.0060.v1>
- García-López, I. M., & Trujillo-Liñán, L. (2025). Generative artificial intelligence in education: ethical challenges, regulatory frameworks and educational quality in a systematic review of the literature. In *Frontiers in Education* (Vol. 10, p. 1565938). Frontiers. <https://doi.org/10.3389/feduc.2025.1565938>
- Gray, S. L., Edsall, D. G., & Parapadakis, D. (2025). AI-based digital cheating at university, and the case for new ethical pedagogies. *Journal of Academic Ethics*, 23(4), 2069. <https://doi.org/10.1007/s10805-025-09642-y>
- Harper, R., Bretag, T., & Rundle, K. (2021). Detecting contract cheating: examining the role of assessment type. *Higher Education Research & Development*, 40(2), 263–278. <https://doi.org/10.1080/07294360.2020.1724899>
- Hertzog, M. (2008). Considerations in determining sample size for pilot studies [Review of considerations in determining sample size for pilot studies]. *Research in Nursing & Health*, 31(2), 180. Wiley. <https://doi.org/10.1002/nur.20247>
- Helm, G., & Hesse, F. W. (2025). Training programmes on writing with AI – but for whom? Identifying students' writer profiles through two-step cluster analysis. *Journal of Writing Research*, 17(1), 1. <https://doi.org/10.17239/jowr-2025.17.01.01>
- Huang, D., Hash, N., Cummings, J. J., & Prena, K. (2025). Academic cheating with generative AI: Exploring a moral extension of the theory of planned behavior. *Computers and Education Artificial Intelligence*, 8, 100424. <https://doi.org/10.1016/j.caeai.2025.100424>
- Hsiao, C.-H., & Tang, K.-Y. (2024). Beyond acceptance: an empirical investigation of technological, ethical, social, and individual determinants of GenAI-supported learning in higher education. *Education and Information Technologies*, 30(8), 10725–10750. <https://doi.org/10.1007/s10639-024-13263-0>
- ICAI. (2020). ICAI | Facts & Statistics. International Center for Academic Integrity. <https://academicintegrity.org/aws/ICAI/pt/sp/facts>
- Johri, A., Hingle, A., & Schleiß, J. (2024). Misconceptions, pragmatism, and value tensions: evaluating students' understanding and perception of generative ai for education. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2410.22289>
- Kangwa, D., Msafiri, M. M., & Fute, A. (2025). Exploring the factors that promote a balance between academic integrity and the effective use of genai tools in higher education: A systematic review. *Journal of Computer Assisted Learning*, 41(5). <https://doi.org/10.1111/jcal.70109>
- Kazu, İ. Y., & Kuvvetli, M. (2023). The Influence of pronunciation education via artificial intelligence technology on vocabulary acquisition in learning English. *International Journal of Psychology and Educational Studies*, 10(2), 480. <https://doi.org/10.52380/ijpes.2023.10.2.1044>
- Kim, J., Klopfer, M., Grohs, J., Eldardiry, H., Weichert, J., Cox, L., & Pike, D. R. (2025). Examining faculty and student perceptions of generative ai in university courses. *Innovative Higher Education*, 50(4), 1281. <https://doi.org/10.1007/s10755-024-09774-w>
- Lancaster, T., & Cotarlan, C. (2021). Contract cheating by STEM students through a file sharing website: a Covid-19 pandemic perspective. *International Journal for Educational Integrity*, 17(1). <https://doi.org/10.1007/s40979-021-00070-0>
- Li, S. (2025). Generative AI and second language writing. *Digital Studies in Language and Literature*. <https://doi.org/10.1515/dsll-2025-0007>
- Leong, W. Y., & Zhang, J. B. (2025). Ethical design of AI for education and learning systems. *ASM Science Journal*, 20(1), 1–9. <https://doi.org/10.32802/asmsci.2025.1917>

- Luo, T., Muljana, P. S., Ren, X., & Young, D. (2024). Exploring instructional designers' utilization and perspectives on generative AI tools: A mixed methods study. *Educational Technology Research and Development*. <https://doi.org/10.1007/s11423-024-10437-y>
- Mahat, D., Neupane, D., & Shrestha, S. (2024). Quantitative research design and sample trends: A systematic examination of emerging paradigms and best practices. *Cognizance Journal of Multidisciplinary Studies*, 4(2), 20. <https://doi.org/10.47760/cognizance.2024.v04i02.002>
- Mahmud, S. (2024). Fostering academic integrity in the age of artificial intelligence. *Advances in educational marketing, administration, and leadership book series*, 1–20. <https://doi.org/10.4018/979-8-3693-0240-8.ch001>
- Mehrpouyan, A. (2025). Ethical boundaries and cybercultural dynamics in virtual drama education: Navigating student-educator interactions and role confusion. *Social Sciences & Humanities Open*, 12, 101712. <https://doi.org/10.1016/j.ssho.2025.101712>
- Mohamed, M. (2024). Exploring ethical dimensions of AI-enhanced language education: A literature perspective. *St-Andrews.ac.uk*. <https://doi.org/318030709>
- Mulenga, R., & Shilongo, H. (2024). Academic integrity in higher education: understanding and addressing plagiarism. *Acta Pedagogica Asiana*, 3(1). <https://doi.org/10.53623/apga.v3i1.337>
- Mursa, R., Patterson, C., McErlean, G., & Halcomb, E. (2025). How many is enough? Justifying sample size in descriptive quantitative research. *Nurse Researcher*, 33(2), 35. <https://doi.org/10.7748/nr.2025.e1958>
- Newton, P. M., & Xiromeriti, M. (2023). ChatGPT performance on MCQ exams in higher education. A pragmatic scoping review. *EdArXiv*, 21.
- Ng, J., Tong, M., Tsang, E. Y. M., Chu, K., & Tsang, W. (2025). Exploring students' perceptions and satisfaction of using genAI-chatgpt tools for learning in higher education: A mixed methods study. *SN Computer Science*, 6(5). <https://doi.org/10.1007/s42979-025-04010-4>
- Pirani, S. A. (2024). Navigating the complexity of sample size determination for robust and reliable results. *International Journal of Multidisciplinary Research & Reviews*, 3(2), 73. <https://doi.org/10.56815/ijmrr.v3i2.2024/73-86>
- Rengifo, M. (2025). Advancing the study of moral disengagement: a summary of current research and future directions. *Edward Elgar Publishing EBooks*, 299–312. <https://doi.org/10.4337/9781035311804.00031>
- Rossi, M., Ciletti, M., Melchiorre, L., & Toto, G. A. (2025). The impact of generative artificial intelligence (genAI) on education: A review of the potential, the risks and the role of immersive technologies. *Education Sciences and Society*, 2, 400–415. <https://doi.org/10.3280/ess2-2024oa18464>
- Salazar, L. R., Peeples, S. F., & Brooks, M. E. (2024). Generative AI ethical considerations and discriminatory biases on diverse students within the classroom. *Advances in educational technologies and instructional design book series*, 191–213. <https://doi.org/10.4018/979-8-3693-0831-8.ch010>
- Sharma, R. C., & Panja, S. K. (2025). Addressing academic dishonesty in higher education: A systematic review of generative AI & rsquo's impact. *Open Praxis*, 17(2). <https://doi.org/10.55982/openpraxis.17.2.820>
- Sevnanarayan, K., & Potter, M.-A. (2024). Generative artificial intelligence in distance education: Transformations, challenges, and impact on academic integrity and student voice. *Journal of Applied Learning & Teaching*, 7(1). <https://doi.org/10.37074/jalt.2024.7.1.41>
- Serwaa, B., Asamoah-Gyawu, J., & Attila, F. L. (2024). Role of self efficacy and personality in academic dishonesty of undergraduate students: Implication for future careers. *Asian journal of advanced research and reports*, 18(7), 10–23. <https://doi.org/10.9734/ajarr/2024/v18i7680>
- Silverman, J. A., Ali, S. A., Rybak, A., Goudoever, van, & Leleiko, N. S. (2023). Generative AI: *Potential and pitfalls in academic publishing*. JPGN Reports, 4(4), e387–e387.
- Siraj, M., Duke, J., & Ploetz, T. (2025). The GenAI generation: Student views of awareness, preparedness, and concern. *arXiv (Cornell University)*. <https://doi.org/10.48550/arxiv.2505.02230>
- Söğüt, S. (2024). Generative artificial intelligence in EFL writing: A pedagogical stance of pre-service teachers and teacher trainers. *Focus on ELT Journal*, 58–73. <https://doi.org/10.14744/felt.6.1.5>
- Stein, G. R., & Wei, Y. (2024). Evaluating student satisfaction: a small private university perspective in Japan. *Higher Education Studies*, 14(2), 27. <https://doi.org/10.5539/hes.v14n2p27>
- Stone, A. (2022). Student perceptions of academic integrity: A qualitative study of understanding, consequences, and impact. *Journal of Academic Ethics*, 21(3), 357–375. <https://doi.org/10.1007/s10805-022-09461-5>
- Tudor, I., Dlab, M. H., Đurović, G., & Horvat, M. (2025). Using clustering techniques to design learner personas for genai prompt engineering and adaptive interventions. *Electronics*, 14(11), 2281. <https://doi.org/10.3390/electronics14112281>
- Uddin, M. M., & Abu, S. (2024). Navigating ethical frameworks to mitigate academic misconduct while leveraging generative AI. *Research Square (Research Square)*. <https://doi.org/10.21203/rs.3.rs-4607113/v1>
- Wilson, J. T. (2023). Gender representation beyond the binary: New possibilities and understandings – *ePrints Soton*. *Soton.ac.uk*. https://eprints.soton.ac.uk/484782/1/Jamie_Wilson_Final_Thesis.pdf
- Wirzal, M. D. H., Nordin, N. A. H. M., Abd, N. S., & Halim, M. (2024). Generative AI in science education: A learning revolution or a threat to academic integrity? A bibliometric analysis. *Journal of Educational Research and Studies: e-Saintika*, 8(3), 319–351.
- Yusuf, A., Pervin, N., & Román-González, M. (2024). Generative AI and the future of higher education: A threat to academic integrity or reformation? Evidence from multicultural perspectives. *International Journal of Educational Technology in Higher Education*, 21(1). <https://doi.org/10.1186/s41239-024-00453-6>
- Zaidy, A. A. (2024). The impact of generative ai on student engagement and ethics in higher education. *journal*

of information technology, cybersecurity, and artificial intelligence., 1(1), 30. <https://doi.org/10.70715/jitcai.2024.v1.i1.004>

Zhang, L., Amos, C., & Pentina, I. (2024). Interplay of rationality and morality in using ChatGPT for academic misconduct.

Behaviour and Information Technology, 44(3), 491. <https://doi.org/10.1080/0144929x.2024.2325023>

Appendix A

Table 4: WSS Values for Optimal Cluster Determination

Number of Clusters (k)	Total Within-Cluster Sum of Squares (WSS)	Rate of Change (Drop in Variance)
1	1332.4	(Baseline)
2	743.8	44.2% (Significant drop)
3	454.1	38.9% (The “Elbow” Point)
4	381.2	16.1% (Marginal gain)
5	344.5	9.6%
6	316.3	8.2%
7	295.1	6.7%
8	278.4	5.7%
9	265.1	4.8%
10	251.9	4.9%

Biographical Statement of Author(s)

Rita R.D. Luise is a serial entrepreneur involved in education and service-based initiatives. She graduated from RMIT (Australia) with a bachelor’s degree and is currently pursuing her postdoctoral studies.

Her main areas of interest are pedagogical design and her sustained engagement with Generative AI has sparked a deep interest in its



transformative possibilities and ethical challenges. Her research examines the use of Generative AI in academic writing support, ethical awareness and value-based decision making among online and distance learners.

Ms. Rita R.D. Luise

GlobalNxt University

Menara HLX

Kuala Lumpur, Malaysia

E-mail: ritaluise2018@gmail.com

Tel: +65 98554515

ORCID ID. <https://orcid.org/0009-0006-9904-547X>