

Voices from the Global South: Myanmar educators' ethical considerations in using generative AI in teaching, learning and assessment

Wai Yan Htut

King Mongkut's University of Technology Thonburi, Bangkok, Thailand

Thu Thu Naing

Cenit College, Naas, Ireland, and

Phyo Wai Tun

University of Warwick, Coventry, UK

Abstract

Purpose – The rapid integration of artificial intelligence (AI) in education presents many opportunities but also raises critical ethical and cultural challenges. This study aims to explore the current application of AI in education in Myanmar and investigates the conflicts between global AI ethical principles and local cultural values.

Design/methodology/approach – Employing a mixed-methods research design, the study utilizes quantitative survey analysis alongside qualitative participatory workshops to provide both breadth and depth to the understanding of educators' perspectives.

Findings – The quantitative findings indicate that while most ethical considerations and practical AI implementation patterns remain consistent across different groups, educators' concern for maintaining human autonomy in AI-assisted education grows significantly with their level of professional experience. Furthermore, the qualitative insights highlight a clear gap between internationally developed (often Global North-centered) AI ethics frameworks and the realities of Myanmar's education context. In particular, socioeconomic challenges and local cultural expectations shape how educators understand and respond to AI. Overall, the study argues that AI ethics in education need to be more locally grounded and responsive to national contexts, rather than relying solely on universal frameworks.

Originality/value – This paper fulfils an identified need to study how ethical AI integration can be enabled in under-resourced and underrepresented educational contexts. By centering the voices of Myanmar educators and documenting their unique perspectives, the study advances knowledge on AI ethics in education and provides practical, context-sensitive guidance for educators, policymakers and researchers.

Keywords Global South, AI ethics, Generative AI, Ethical guidelines, Myanmar education, Teacher perspectives

Paper type Research article

1. Introduction

The rapid advancement of generative artificial intelligence (GenAI) has transformed educational practices, offering new opportunities for teaching, learning and assessment. GenAI tools are now widely used by educators in creating lesson content and personalizing instruction, but their integration also raises significant ethical concerns, including transparency, data privacy and the shifting roles of teachers and students (Moorhouse *et al.*, 2023). Consequently, there has been a sudden proliferation of ethical frameworks aimed at governing the responsible use of AI globally (e.g. Floridi, 2023). However, the global



discourse remains largely dominated by perspectives, priorities and frameworks originating from the Global North, and those frameworks may not adequately address the unique challenges faced by educators in the Global South (Adams, 2021; Nemorin, 2024; Roche *et al.*, 2021).

There remains a critical gap in understanding how universal AI application principles conflict with local cultural norms, infrastructure limitations and educational practices in developing nations. Global South contexts, such as Myanmar, are often characterized by severe infrastructural limitations, socioeconomic disparities and deeply rooted cultural traditions that shape educational practices and ethical priorities differently than in high-income countries (Anderson and Mahapatra, 2024; Lall and South, 2014; Vijayakumar, 2024). While global AI ethics frameworks frequently emphasize principles like transparency and autonomy, for instance, they tend to overlook the situated, local realities of the Global South – a systemic marginalization linked to the coloniality of power in data epistemologies (Ricaurte, 2019). Because AI research, development and policy-making are heavily concentrated in wealthier nations, ethical templates are frequently imposed on Global South contexts without sufficient adaptation to local needs, values or vulnerabilities (Vijayakumar, 2024).

This study aims to address this critical gap in understanding context-specific ethical considerations by amplifying the voices of Myanmar educators. Using a mixed-methods approach, it aims to explore the teaching and learning principles that guide educators' GenAI use, the ethical considerations by educators and their alignment with local contexts. The findings aim to inform the development of practitioner-informed, contextually relevant ethical guidelines for GenAI integration in education in Myanmar and the Global South.

This study provides empirical evidence to challenge the hegemony of globalized AI ethics by amplifying local perspectives and explicitly tying them to educational outcomes. By identifying the root causes of the friction between global AI mandates and local educational cultures, this research contributes to the important process of decolonizing the ethics of AI in education (Nemorin, 2024). It moves beyond abstract principles to offer insights for policymakers, institutions and developers to construct context-sensitive, ethically aligned AI systems that genuinely prioritize regional human wellbeing across diverse cultural landscapes (Shahriari and Shahriari, 2017).

2. Literature review

The field of AI ethics has experienced rapid expansion, with Floridi (2023) documenting extensive proliferation of ethical frameworks by the early 2020s, reaching over 160 distinct sets of principles according to Algorithm Watch's AI Ethics Guidelines Global Inventory. This growth initially generated concerns about fragmentation and potential confusion among stakeholders seeking guidance for responsible AI implementation, particularly regarding GenAI applications in educational contexts.

However, scholarly analysis has revealed unexpected coherence across these seemingly diverse approaches. Floridi's (2023) systematic comparison of prominent ethical frameworks – encompassing initiatives from the Asilomar AI Principles (Future of Life Institute, 2017), Montreal Declaration for Responsible AI (Université de Montréal, 2017), IEEE's Ethically Aligned Design (Shahriari and Shahriari, 2017), European Group on Ethics statement (European Group on Ethics in Science and New Technologies, 2018), UK House of Lords recommendations (House of Lords Artificial Intelligence Committee, 2018) and Partnership on AI tenets (Partnership on AI, 2018, as cited in Floridi and Cows, 2019) – demonstrates substantial alignment across 47 individual principles despite originating from different geographical, institutional and stakeholder contexts.

Building on this foundational understanding, Funa and Gabay's (2025) comprehensive meta-synthesis of AI policies and guidelines from 2020 to 2024 provides crucial insights into how these theoretical frameworks translate into practical implementation in educational contexts. Their analysis of 23 sources across diverse global contexts reveals

that while ethical principles demonstrate remarkable consistency, their implementation faces significant practical constraints that vary considerably between developed nations and Global South contexts.

2.1 Theoretical foundation

Contemporary AI ethics frameworks draw extensively from established bioethical principles, reflecting a conceptual approach that views AI as a new form of agency rather than mere computational capability. This theoretical positioning makes bioethics particularly relevant as an analogous domain for addressing ethical challenges in GenAI deployment in educational settings (Floridi, 2013). The four foundational principles of bioethics—beneficence, nonmaleficence, autonomy and justice – provide a robust starting point for GenAI governance in education (Beauchamp and Childress, 2013).

Promoting beneficial outcomes (beneficence) appears as the most universally recognized principle across AI ethics frameworks. Various formulations focus on promoting well-being, preserving human dignity and ensuring environmental sustainability. The Montreal Declaration calls for AI development to “promote the well-being of all sentient creatures,” while the IEEE principles prioritize “human well-being as an outcome in all system designs” (Floridi *et al.*, 2018, p. 696). In the context of GenAI in education, this principle encompasses using these technologies to enhance learning outcomes, support pedagogical goals and promote student welfare. Funa and Gabay’s (2025) meta-synthesis reinforces this finding, noting that beneficence consistently emerges across diverse cultural and institutional contexts, though its interpretation may vary according to local educational priorities and resource constraints.

Preventing harm (nonmaleficence) addresses the imperative to prevent harm, encompassing both privacy protection and broader security concerns. The frameworks particularly emphasize preventing infringements on personal privacy and avoiding the misuse of AI technologies, though they leave ambiguous whether the responsibility lies with AI developers or the systems themselves (Floridi, 2023). In GenAI applications, this principle is particularly relevant given concerns about academic integrity, data privacy and the potential for these tools to undermine critical thinking skills. Funa and Gabay (2025) identify data privacy as a central concern across their analyzed sources, emphasizing that “educational institutions must implement robust data protection measures to safeguard sensitive information collected through AI tools” (p. 11), highlighting the practical urgency of this principle in educational contexts.

Preserving human agency (autonomy) presents complex considerations in the GenAI context. Unlike traditional bioethics where autonomy is typically compromised involuntarily, GenAI adoption involves willingly ceding decision-making power to artificial agents. The frameworks emphasize the need for “meta-autonomy” – maintaining human control over which decisions to delegate and preserving the ability to override automated systems (Floridi *et al.*, 2018, p. 698). This concept gains additional significance in educational contexts where GenAI tools like ChatGPT can generate content that might replace student thinking, where Funa and Gabay (2025) note that maintaining human oversight ensures “AI technologies enhance rather than replace the roles of teachers and students” (p. 6).

Ensuring fair distribution (justice) encompasses multiple dimensions, including using AI to correct past wrongs such as eliminating discrimination, ensuring equitable distribution of AI benefits and preventing the creation of new harms or undermining of existing social structures (Floridi, 2023). In GenAI applications, this principle addresses concerns about differential access to these tools and their potential to exacerbate educational inequalities. Funa and Gabay’s (2025) analysis reveals that justice considerations are particularly acute in Global South contexts, where “disparities in access to AI tools and resources can exacerbate existing inequalities in education” (p. 10).

The fifth principle (explicability) incorporates both epistemological intelligibility (answering “how does it work?”) and ethical accountability (answering “who is responsible

for the way it works?”) (Floridi and Cowls, 2019). In the context of GenAI tools, explicability becomes particularly complex given the “black box” nature of many large language models and the difficulty in understanding how these systems generate specific outputs.

Funa and Gabay’s (2025) meta-synthesis provides empirical support for this principle’s importance, identifying “transparency and privacy concerns” as major practical constraints in AI implementation. They note that transparency issues arise when “the decision-making processes of AI systems are not clearly communicated, leading to uncertainty and skepticism about the technology’s role in education” (p. 11). For GenAI applications in education, this principle encompasses the need for educators and students to understand how these tools generate responses, their limitations and the appropriate contexts for their use.

While these ethical principles demonstrate great consistency in theory, their practical implementation faces many constraints. A critical problem with the current global landscape is that the global discourse on AI ethics, including the authorship of these very guidelines, remain largely dominated by perspectives, priorities and frameworks originating from the Global North (Adams, 2021; Nemorin, 2024; Roche *et al.*, 2021). This sets the stage for an epistemic imbalance, as universalist ethical templates are frequently promoted without addressing the severe socioeconomic disparities or unique cultural realities of the developing world.

2.2 Pedagogical foundations: AI and learning theories

Before applying global ethical principles, the integration of AI in education (especially, teaching, learning and assessment), it is important to ground it in established learning theories, as effectiveness of AI depends largely on pedagogical rationale (Elstad, 2024). The evolution of educational theory reflects a growing understanding of learning as a complex, multifaceted process that any single theoretical framework cannot fully explain. Understanding how generative AI intersects with these foundational theories is essential for effective pedagogical implementation (Moussa, 2024).

2.2.1 AI and behaviorist approaches. Behaviorist learning theory, developed through the pioneering work of Skinner (1938) and Watson (1998), focuses on observable behavior changes as the primary indicators of learning, emphasizing stimulus-response relationships and the role of reinforcement in shaping learning outcomes. AI technologies support behaviorist learning approaches through several mechanisms. They facilitate drill and practice exercises, allowing students to receive consistent repetition of key concepts necessary for skill development (Bastani *et al.*, 2024). These tools provide immediate feedback on student performance, enabling real-time correction and reinforcement (Akila *et al.*, 2023). According to Elstad (2024), AI systems enable reinforcement learning while simultaneously requiring students to develop self-discipline, as the technology can track progress but cannot enforce persistence without student commitment.

2.2.2 AI and cognitive approaches. Cognitive learning theory represents a significant shift from behaviorist approaches by focusing on internal mental processes, schema development and the active role of learners in constructing understanding (Piaget, 2007; Bruner, 1966; Ausubel, 1968). From a cognitive learning perspective, AI offers several advantages. Klopfer *et al.* (2024) demonstrate that AI enhances information processing without creating pressure. AI functions can also support problem-solving processes while allowing students to maintain agency in their learning journey (Koh *et al.*, 2023). Notably, AI can both challenge and support critical thinking development while providing scaffolded assistance when needed (Hodges and Kirschner, 2024).

2.2.3 AI and constructivist approaches. Constructive learning theory, rooted in the work of Vygotsky (1978), Dewey (1997) and von Glasersfeld (1995), emphasizes the active construction of knowledge through experience and social interaction, highlighting the necessity of active exploration rather than the passive reception of information. AI demonstrates significant potential within constructivist learning frameworks. Elstad (2024) highlights how AI facilitates active learning processes by providing personalized pathways

through content. For institutions with resource constraints, AI enables more student-centered learning by addressing limitations in teaching staff and materials (Akila *et al.*, 2023). Personalized learning experiences supported by AI allow students to build knowledge structures that connect meaningfully to their existing understanding and interests (Elstad, 2024). Mollick and Mollick (2023) demonstrate that AI can provide multiple examples and explanations of concepts, supporting the constructivist principle that learners benefit from varied knowledge representations.

2.2.4 AI and social learning approaches. Social learning theory, developed by Bandura (1977) and extended through communities of practice research by Lave and Wenger (1991, Wenger, 1998), bridges individual cognitive processes and social environmental factors to explain how learning occurs within social contexts through observation, modeling and social interaction. AI technologies increasingly support social learning theories in educational contexts. They facilitate collaborative learning environments by providing platforms for shared knowledge construction and group problem-solving (Koh *et al.*, 2023). According to Hodges and Kirschner (2024), AI enables more sophisticated peer interaction and review processes by providing additional perspectives and feedback that complement human contributions. Chen *et al.* (2023) demonstrate that AI can serve as a knowledge-building partner in collaborative learning environments.

In the field of educational evaluation, traditional assessment methods are being fundamentally challenged by AI technologies that can generate essays, solve math problems, write code and even pass standardized tests with impressive results. Research indicates that AI systems like ChatGPT can produce academic work that qualifies for high marks in university settings, prompting educators to rethink assessment strategies across disciplines (Terwiesch, 2023). Multiple studies reveal that AI tools allow students to generate original essays almost instantly with little cost, making it increasingly difficult to distinguish machine-generated from student-generated answers, which creates significant concerns about academic integrity (Kasneci *et al.*, 2023).

Conversely, AI is particularly valuable for formative assessment, which is the ongoing evaluation that informs teaching strategies. Tools like AI tutors are designed to guide students without giving direct answers, encouraging deeper learning through Socratic questioning techniques (Mollick and Mollick, 2023). These technologies can help identify knowledge gaps and learning obstacles in real-time, allowing for immediate intervention and support. With advancements in natural language processing, these systems are increasingly capable of evaluating written responses and providing feedback on writing skills (Zawacki-Richter *et al.*, 2019).

Despite these promising opportunities, integrating AI with these learning approaches presents significant pedagogical challenges, as Bastani *et al.* (2024, p. 1) identified a primary concern that students using AI as a “crutch” can experience diminished learning outcomes when they rely on AI-generated solutions without engaging with productive cognitive effort. Hodges and Kirschner (2024) further caution that AI tools can deprive students of valuable learning opportunities if they bypass the desirable difficulties that actively promote deep learning.

In addition, Implementing AI in education must account for diverse cultural contexts that significantly impact its effectiveness and appropriateness. Chen *et al.* (2023) emphasize that AI systems influence students’ engagement with particular cultural factors. Language accessibility presents another crucial consideration, as Elstad (2024) highlights that AI tools may be biased against non-native English writers.

2.3 Global south perspectives in AI ethics

2.3.1 Under-representation of the global south in ethical considerations for AI in education.

The global discourse on AI ethics remains largely dominated by perspectives, priorities and frameworks from the Global North (Adams, 2021; Roche *et al.*, 2021; Nemorin, 2024). This

dominance has resulted in a significant gap in understanding how AI technologies impact educational contexts in the Global South, where infrastructural, socioeconomic and cultural realities differ markedly from those in high-income countries (Vijayakumar, 2024; Ricaurte, 2019).

AI ethics frameworks – such as those developed by the OECD and UNESCO – tend to reflect the regulatory, cultural and technological conditions of the Global North (Vijayakumar, 2024). These frameworks often prioritize issues, like privacy, transparency and technical trustworthiness, but they do not fully address the social costs, power imbalances or local realities experienced in the Global South (Adams, 2021; Vijayakumar, 2024). Because AI research, development and policy-making are heavily concentrated in wealthier nations, the voices and experiences of educators, policymakers and communities from the Global south are systematically marginalized (Adams, 2021; Ricaurte, 2019). This imbalance results in a form of technological and epistemic colonialism, where foreign knowledge production, technologies and ethical templates are imposed on developing contexts without sufficient adaptation to local needs, values or vulnerabilities (Ricaurte, 2019; Vijayakumar, 2024).

2.3.2 Unique challenges and contexts of the global south. Countries in the Global South face a distinct set of challenges that shape their ethical considerations around AI in education. To fully understand this phenomenon, it is essential to look beyond the context of a single nation and examine how parallel struggles across the broader Global South involve infrastructural limitations, digital divides and algorithmic bias that risk exacerbating existing inequalities (Ricaurte, 2019; Vijayakumar, 2024). Firstly, there are strong anti-colonial concerns about the exploitation of data and resources, with fears that AI will perpetuate dependency on the Global North (Vijayakumar, 2024). Secondly, Algorithmic bias and digital divides risk exacerbating existing inequalities, particularly for marginalized and vulnerable groups (Vijayakumar, 2024; Ricaurte, 2019). In addition, AI systems often reflect dominant sociocultural norms, which can marginalize local languages, values and practices (Vijayakumar, 2024). Moreover, the proliferation of “soft law” instruments professing ethical principles often lacks effective accountability mechanisms, leaving the Global South with limited means to hold powerful actors accountable (Vijayakumar, 2024).

The impact of these challenges is clearly visible across diverse regional contexts. In Latin America, for example, the integration of AI is heavily impacted by the “coloniality of knowledge,” a dynamic where structural inequalities restrict equitable access to education and populations suffer the imposition of foreign technologies that reflect the cultural optimization ideals of the Global North (Mancilla-Caceres and Estrada-Villalta, 2022, p. 2). Furthermore, severe socioeconomic disparities mean that lower-income students often rely primarily on smartphones and free, limited AI tools for learning, which exacerbates the digital divide within the region’s higher education systems (Valdivieso and González, 2025). Similarly, in Sub-Saharan Africa, the implementation of critical digital pedagogy and effective AI adoption is severely hindered by infrastructural challenges, such as unreliable electricity, poor Internet connectivity and low digital literacy among educators (Ncube and Tawanda, 2025).

These challenges are compounded by infrastructural limitations (such as unreliable electricity and Internet access), socioeconomic disparities and a lack of local capacity for AI governance (Vijayakumar, 2024; Ricaurte, 2019). As a result, the ethical implications of AI in education are deeply context-dependent and cannot be fully addressed by top-down universalist frameworks (Adams, 2021; Monasterio Astobiza *et al.*, 2022). Consequently, there is growing recognition of the need for more inclusive, context-sensitive approaches to AI ethics in education. This requires active participation from Global South educators, policymakers and communities in the development of ethical guidelines, as well as the creation of frameworks that reflect local realities, priorities and values (Vijayakumar, 2024; Ricaurte, 2019). Grassroots initiatives – such as the development of local language datasets and community-driven governance models – demonstrate the potential for more equitable and effective approaches to AI ethics (Vijayakumar, 2024; UNESCO, 2024).

2.4 Myanmar as a representative context

Myanmar exemplifies how Global South realities shape and challenge the application of AI ethics in education. The country's unique infrastructural, cultural and political context reveals the limitations of universalist frameworks and underscores the need for locally grounded ethical considerations.

2.4.1 Infrastructural realities. Myanmar's technological landscape is marked by significant disparities: rural electrification remains low and mobile Internet penetration is limited, especially in conflict-affected areas (Tun and Smith, 2026). Teachers often face connectivity-dependent ethical dilemmas, such as relying on outdated AI-generated content during Internet blackouts (Borenstein and Howard, 2021). These infrastructural challenges exacerbate algorithmic bias, as AI tools trained on limited or non-representative data may misrepresent ethnic minority histories or local realities (Borenstein and Howard, 2021).

2.4.2 Cultural and contextual differences. Because dominant AI ethics frameworks are rooted in Western epistemologies, they place a heavy emphasis on principles like transparency and autonomy, advocating for the explicit disclosure of AI usage in classrooms (Floridi and Cows, 2019). However, cultural norms in Myanmar might shape how educators perceive ethical principles like transparency and autonomy. In Myanmar, teachers are regarded as the custodians of knowledge and moral authority. This concept is deeply embedded in Myanmar culture and shapes classroom dynamics, teacher–student relationships and perceptions of authority (Seekins, 2017; Lall and South, 2014).

For example, the *gadaw* ceremony is a formal act of paying homage to teachers, elders and monks, reflecting hierarchical social structures and respect for authority. This ceremony is commonly practiced in schools and is an important part of Myanmar's educational culture (Lall and South, 2014; Seekins, 2017). Consequently, while Western frameworks emphasize transparency, requiring educators to disclose AI use in lesson planning or assessment, many Myanmar teachers may view AI tools as extensions of their pedagogical authority. As a result, they may choose not to inform students about the use of AI in content creation, believing that disclosure could undermine their professional status or disrupt established classroom hierarchies.

2.5 Challenges and research gaps in AI integration

Integrating AI with various learning approaches presents significant challenges and promising opportunities (Hodges and Kirschner, 2024). AI offers substantial opportunities to enhance education; for instance, Mollick and Mollick (2023) demonstrate that AI can support multiple teaching strategies simultaneously, providing personalized explanations, examples and feedback that accommodate diverse learning preferences. However, Bastani *et al.* (2024) identified a primary concern that students using AI as a “crutch” can experience diminished learning outcomes, particularly when they rely on AI-generated solutions without engaging in productive cognitive effort. Hodges and Kirschner (2024) further caution that AI tools can deprive students of valuable learning opportunities if they bypass the desirable difficulties that promote deep learning.

Despite these challenges, attempting to mitigate them through restrictive measures creates significant equity issues. As Klopfer *et al.* (2024) noted, school policies that block AI tools may inadvertently disadvantage lower-income students who rely solely on school-issued devices. Moreover, AI assessment systems often collect substantial student data, raising concerns about privacy and security (Xia *et al.*, 2024). Educational institutions must carefully consider what information is being gathered, how it is stored and who has access to it (Xia *et al.*, 2024). The ethical deployment of AI requires transparent policies about data collection and use to protect student information.

While research has documented how AI can support various learning theories and assessment processes, there remains limited empirical evidence of how educators in the Global South actually navigate these complex pedagogical and ethical decisions when implementing

GenAI tools in their classes. To address this epistemic gap, counteract the under-representation of the Global South in AI discourse and broaden the understanding of context-specific ethical considerations, this research aims to answer four questions about GenAI integration in the Global South, specifically the Myanmar educational context.

- (1) What teaching and learning principles do educators in Myanmar base on when they use GenAI in their teaching, learning and assessment?
- (2) What ethical considerations do they consider when they use GenAI?
- (3) Do these ethical considerations align with the unified framework of ethical AI?
- (4) What ethical considerations do they recommend when integrating GenAI into the classroom?

3. Methodology

3.1 Research design: a mixed-methods approach

To investigate the ethical considerations of Myanmar educators regarding Generative AI in teaching, learning and assessment, this study employed a mixed-methods research design (Creswell, 2014). Educational environments are highly complex and multi-layered, meaning that relying on a single methodological paradigm often yields an incomplete picture of classroom realities (Mejeh *et al.*, 2023). This stance was particularly appropriate for this study where we investigate ethical considerations in AI use, where multiple perspectives and contextual factors play crucial roles in shaping educators' views (Mejeh *et al.*, 2023). A mixed-methods approach was specifically chosen because it embraces methodological pluralism, allowing researchers to combine the complementary strengths of both quantitative and qualitative techniques while minimizing their respective, non-overlapping weaknesses (Johnson and Onwuegbuzie, 2004; Moss and Haertel, 2016). Rooted in the philosophical framework of pragmatism, this approach prioritizes the research questions over paradigm purity, utilizing whichever combination of methods works best to provide a superior, more comprehensive understanding of the phenomenon (Johnson and Onwuegbuzie, 2004).

We specifically adopted a sequential mixed-methods design, which involves collecting and analyzing quantitative data first, and then building on those results to explain them in more detail through qualitative research (Creswell, 2014). We began with a quantitative survey followed by qualitative participatory workshops. This sequential approach allowed us to first identify broad patterns and trends through the survey and then explore these findings in greater depth through interactive workshops. The integration of these methods provided both breadth and depth to our understanding of Myanmar educators' perspectives on ethical AI use in education.

The design consisted of two interconnected phases:

Phase 1 (Quantitative–priority phase):

A survey of 111 Myanmar educators provided breadth of understanding by identifying patterns in AI implementation practices and ethical considerations across diverse educational contexts. This phase addressed the research questions regarding what ethical considerations educators prioritize and how these align with established frameworks. A primary justification for starting with quantitative methods is its strength in identifying broad trends and relationships across diverse educational contexts (Johnson and Onwuegbuzie, 2004).

Phase 2 (Qualitative–explanatory phase):

While quantitative data can effectively map *what* ethical considerations are prevalent among Myanmar educators, it is often too abstract to explain *why* these trends exist in specific local contexts. To resolve this, the second phase utilized qualitative participatory workshops to

probe the initial survey results in depth. Participatory workshops with six educators provided depth of understanding by exploring the “why” and “how” behind quantitative patterns. This phase illuminated contextual factors, cultural influences and nuanced perspectives that numbers alone could not capture ([Johnson and Onwuegbuzie, 2004](#)).

Overall, the main justification for this mixed-methods approach is that it bridges the gap between macro-level theoretical frameworks, such as the Unified Framework of Ethical AI use, and the micro-level (classroom) realities experienced by educators in the Global South ([Mejeh et al., 2023](#)). By first using a survey to generate broad statistical patterns, and then using qualitative tools to capture the nuanced cultural and infrastructural frictions in Myanmar, the resulting integration produces a more robust and validated answer to the research questions than monomethod research could achieve ([Creswell, 2014](#); [Johnson and Onwuegbuzie, 2004](#)).

3.2 Theoretical framework

The study was based on the Unified Framework of Five Principles for AI in Society ([Floridi and Cowls, 2019](#)). This framework served as a theoretical lens through which we examined Myanmar educators’ ethical considerations. However, recognizing the potential limitations of applying Global North frameworks to Global South contexts, we employed participatory methods to facilitate the emergence of contextually relevant ethical considerations.

3.3 Participants and sampling

3.3.1 Survey participants. For the quantitative phase of the study, 111 Myanmar educators participated in the survey. Participants were recruited through a combination of snowball sampling and social media outreach to ensure diverse representation across different educational contexts within Myanmar. This approach was chosen to maximize reach within a challenging research environment where traditional random sampling might be impractical. All participants were currently teaching in Myanmar educational institutions and had some experience with or interest in using AI tools in their practice ([Cumbo and Selwyn, 2021](#)).

3.3.2 Workshop participants. From the survey respondents who indicated interest in participating in follow-up activities, we invited educators to join the participatory workshop. Initially, 25 participants expressed interest, but ultimately, six educators attended the workshop. The planned second session was cancelled due to a devastating earthquake that affected the region and the participants. While the smaller number of workshop participants represents a limitation, it aligns with recommendations for smaller groups in participatory workshops to ensure meaningful engagement and in-depth discussion ([Creswell, 2014](#)).

3.4 Data collection methods

3.4.1 Questionnaire. A comprehensive questionnaire was developed to collect data from 111 Myanmar educators. The questionnaire was structured into four key sections:

- (1) AI in teaching and learning
- (2) AI in assessment
- (3) AI in educational theories
- (4) Ethical considerations using the unified framework

The instrument employed a 5-point Likert scale (1–5) to measure participants’ levels of agreement with various statements related to ethical consideration of AI in teaching, learning and assessment. The questionnaire was followed by an open-ended “Why?” question to capture qualitative insights into participants’ reasoning. This combination of closed and open-ended questions exemplifies the integration of quantitative and qualitative approaches within a mixed-methods design ([Mejeh et al., 2023](#)).

The questionnaire was collaboratively created by three researchers, all with contextual knowledge of Myanmar's educational system. The items related to ethical considerations were based on the Unified Framework of Five Principles for AI in Society (Floridi and Cowsls, 2019), ensuring alignment with established theoretical perspectives while allowing for contextual adaptation through open-ended responses.

3.4.2 Participatory workshop. The participatory workshop was designed as a bottom-up approach to complement the survey data and allow educators to elaborate on their questionnaire responses. This method aligns with participatory research principles that aim to empower participants as agents of knowledge creation rather than mere study subjects (Cumbo and Selwyn, 2021).

The two-hour workshop followed a structured format:

- (1) Introduction and aim statement (5 min)
- (2) Warm-up poll/discussion on ethical implications of AI in education (5 min)
- (3) Demonstration of an AI tool (5 min)
- (4) Discussion of ethical implications (10 min), introducing Global North/South perspectives
- (5) Breakout room discussion: "What ethical considerations should we consider in our context? Global south?" (30 min, with Padlet for note-taking)
- (6) Introduction to the Unified Framework (10 min)
- (7) Breakout room discussion: "What ethical considerations should we consider in our context with reference to the framework?" (30 min, using Google Forms)
- (8) Conclusion and introduction to follow-up interviews (5 min)

Two researchers facilitated each breakout room to ensure balanced discussions and comprehensive documentation. The workshop was conducted in Burmese, the participants' native language, to enable them to express their thoughts and ideas more freely and naturally. This linguistic choice was critical for capturing nuanced perspectives on ethical considerations that might be lost in translation during the workshop itself.

3.4.3 Instrument reliability, validity and translation validity. 3.4.3.1 Reliability. The questionnaire's internal consistency was assessed using Cronbach's alpha coefficient for each of the three primary domains. As shown in Table 1, all domains demonstrated excellent reliability, well exceeding the acceptable threshold of $\alpha = 0.70$ (Nunnally, 1978).

The AI in Educational Theories and Ethical Considerations domain incorporated items based on Floridi and Cowsls' (2019) Unified Framework of Five Principles for AI in Society (beneficence, non-maleficence, autonomy, justice and explicability). While these principles were analyzed separately in the findings section to provide nuanced understanding of educators' ethical considerations, the combined domain demonstrated excellent internal consistency ($\alpha = 0.921$), indicating that items reliably measured educators' ethical perspectives on AI use in education.

Table 1. Reliability statistics for questionnaire domains

Domain	Number of items	Cronbach's alpha	Interpretation
AI in teaching and learning	22	0.963	Excellent
AI in assessment	16	0.954	Excellent
AI in educational theories and ethical considerations	20	0.921	Excellent
Overall instrument	58	0.974	Excellent

The exceptionally high overall reliability ($\alpha = 0.974$) demonstrates strong internal consistency across the entire instrument, suggesting that the questionnaire items consistently measured the intended constructs related to AI use in Myanmar education.

3.4.3.2 Validity. Content validity was established through multiple procedures:

- (1) *Theoretical grounding*: Items addressing ethical considerations were systematically derived from [Floridi and Cows' \(2019\)](#) Unified Framework of Five Principles for AI in Society, ensuring alignment with established ethical AI principles recognized in international discourse.
- (2) *Expert development*: The questionnaire was collaboratively constructed by three researchers with extensive knowledge of Myanmar's educational context, ensuring cultural relevance and contextual appropriateness for the target population.
- (3) *Contextual adaptation*: Items were developed to reflect Myanmar's specific educational realities (including infrastructure challenges, cultural norms and institutional diversity) while maintaining theoretical fidelity to the underlying frameworks of both AI implementation and ethical considerations.
- (4) *Framework integration*: Items addressing AI in teaching, learning and assessment were informed by established learning theories (behaviorist, cognitive, constructivist and social learning approaches) as discussed in the literature review, ensuring pedagogical validity.

3.4.3.3 Translation validity. The questionnaire was developed in English, translated into Burmese by bilingual members of the research team and reviewed for conceptual equivalence to ensure accurate meaning transfer between languages. The workshop's use of Burmese further validated that participants understood and could meaningfully engage with the concepts measured in the questionnaire.

3.4.4 *Qualitative reliability, validity and controlling bias*. Because qualitative research involves interpretative analysis and relies on the researcher as the primary instrument for data collection, it is crucial to embed safeguards that ensure the trustworthiness of the findings ([Creswell, 2014](#)). To directly address the inherent subjectivity of qualitative inquiry and ensure the rigorous evaluation of our participatory workshop data, we implemented systematic strategies to control our own researcher bias, establish qualitative validity and ensure qualitative reliability.

3.4.4.1 Controlling researcher bias. In qualitative research, we must reflect on how our "personal backgrounds, cultures, and experiences hold the potential to shape [our] their interpretations" of the data ([Creswell, 2014](#), p. 186). To mitigate personal subjectivity, we actively engaged in reflexivity, continually clarifying how our "insider" knowledge of the Myanmar educational system might influence our analysis. Furthermore, to restrict the imposition of ungrounded biases, we anchored our initial qualitative coding process in an *a priori* deductive codebook based on the five principles of the Unified Framework of Ethical AI ([Floridi and Cows, 2019](#)). This theoretical anchoring ensured our analysis remained grounded in established scientific parameters, while we still permitted new inductive codes to emerge organically from the participants' lived realities.

3.4.4.2 Qualitative validity. Qualitative validity refers to the procedures we used to check whether our findings accurately reflect the realities and standpoints of the participants ([Creswell, 2014](#)). To establish this, we employed data triangulation by constantly comparing the qualitative workshop insights against our quantitative survey data, using converging evidence from both phases to build a coherent justification for the established themes. We also strengthened validity through the provision of rich, thick descriptions of the educators' contexts and by intentionally presenting discrepant information, such as differing opinions among teachers regarding data security, to ensure our final account was realistic and multifaceted ([Creswell, 2014](#)).

More importantly, we strictly maintained translation validity throughout the data collection process. We conducted the participatory workshops entirely in Burmese, allowing the educators to express nuanced cultural concepts freely and naturally. To preserve these cultural and linguistic nuances, our bilingual research team members carefully translated representative quotes into English. We then rigorously cross-checked these translations with other bilingual team members for conceptual equivalence to ensure no meaning was lost in translation.

3.4.4.3 Qualitative reliability. Qualitative reliability indicates that our analytical approach is stable and consistent across different researchers and throughout the duration of the project (Creswell, 2014). To ensure consistency at the data preparation level, we transcribed all workshop recordings verbatim in Burmese and meticulously cross-checked them against the original audio files to eliminate transcription errors.

During the analysis phase, we strictly enforced reliability through intercoder agreement protocols (Creswell, 2014). Rather than relying on a single coder, multiple researchers on our team conducted comprehensive independent coding. Following this independent phase, we held comparison and calibration meetings where we discussed our coding decisions and reached a consensus to resolve any discrepancies. This collaborative process prevented “code drift” (a shift in the meaning of the codes during the process) and ensured we applied the ethical themes consistently across all text segments (Creswell, 2014).

3.5 Data analysis

3.5.1 Quantitative analysis. The quantitative data from the questionnaire were analyzed using SPSS 25 software. Descriptive statistics, including means, standard deviations and frequency distributions, provided a broad overview of respondents’ stated practices and ethical considerations in using AI in teaching, learning and assessment. To examine whether differences between groups were statistically significant, inferential statistical tests were conducted. Mann–Whitney *U* tests were used to compare composite domain scores and ethical sub-domain scores between public- and private-sector educators, the two largest institutional groups in the sample ($n = 28$ and $n = 75$, respectively). Kruskal–Wallis *H* tests were used to compare scores across four experience-level groups (0–2 years, $n = 12$; 3–5 years, $n = 34$; 6–10 years, $n = 43$; more than 10 years, $n = 22$). Non-parametric tests were selected because the data were collected on ordinal Likert scales and could not be assumed to follow normal distributions (Field, 2018). For the Kruskal–Wallis test, Dunn’s post-hoc pairwise comparisons with Bonferroni adjustment were conducted where significant differences were found. Effect sizes were calculated using $r = Z/\sqrt{N}$ for Mann–Whitney *U* tests and $\eta^2 = (H - k + 1)/(N - k)$ for Kruskal–Wallis tests. The quantitative results helped identify patterns and trends that informed the subsequent qualitative exploration.

3.5.2 Qualitative data analysis. The qualitative component of this study comprised two data sources: (1) open-ended survey responses from 77 participants and (2) audio-recorded workshop discussions from 6 participants. Both sources were analyzed using a hybrid deductive-inductive thematic analysis approach.

Data preparation. Workshop recordings were transcribed verbatim in Burmese by members of the research team. Transcripts were verified for accuracy through cross-checking with other researchers from the participating institutions who reviewed the recordings against the transcriptions to ensure completeness and fidelity to the original discussions. The 77 open-ended survey responses, also provided in Burmese, were compiled into a single document for analysis. All qualitative data remained in Burmese during the initial coding process to preserve linguistic and cultural nuances.

Analytical framework. A hybrid deductive-inductive approach was employed (Fereday and Muir-Cochrane, 2006), combining:

- (1) *Deductive analysis:* Using predetermined codes from Floridi and Cows’ (2019) Unified Framework of Five Principles for AI in Society

AIIE
2,4

(2) *Inductive analysis:* Identifying emergent themes specific to Myanmar’s educational and cultural context that were not captured by the existing framework

This approach allowed the research team to systematically examine alignment with established ethical principles in Global North contexts.

Initial coding framework. Five deductive codes were established *a priori* based on the unified framework:

- 76
- (1) Beneficence

(2) Non-maleficence

(3) Autonomy

(4) Justice

(5) Explicability

Coding procedure and inter-coder reliability: The coding process followed a systematic approach: Initial Independent Coding, Comparison and Calibration, Comprehensive Independent Coding, Inductive Code Development and Consensus Building and Validation. The researchers approached independent coding followed by discussion provided reliability checking. The requirement that all coding decisions be discussed and agreed upon by all team members ensured consistency in code application.

Translation procedures: Representative quotes selected for inclusion in the findings were translated from Burmese to English by bilingual members of the research team. To ensure conceptual equivalence and meaning preservation, translations were reviewed by the other bilingual team members, and cultural context and linguistic nuances were preserved.

3.5.3 Mixed-methods integration. The qualitative analysis was designed to explain and expand upon quantitative patterns identified in the survey data. Integration occurred during the workshop design and during the qualitative analysis process.

3.6 Methodological limitations

Several methodological limitations should be acknowledged. While the survey sample ($n = 111$) was adequate for descriptive analysis, the workshop phase faced substantial constraints. Only 6 of 25 interested participants attended, and the planned second session was cancelled due to a regional earthquake, limiting the depth of qualitative data. Data collection during Myanmar’s period of political instability and the online format may have affected participants’ openness and discussion depth.

4. Findings

4.1 Teaching learning principles and use of AI

4.1.1 Frequency of AI tool usage. The data reveal varying degrees of GenAI adoption among Myanmar educators. Using a 5-point Likert scale (where higher values indicate greater frequency), 21.6% of educators report the highest frequency of use (5.00), while 26.1% report high usage (4.00). The largest group (32.4%) falls in the middle range (3.00), with smaller percentages reporting low usage (10.8% at 2.00) and very low usage (9.0% at 1.00). These findings suggest that the majority of educators (80.1%) use AI tools at moderate to high frequencies (see Figure 1).

4.1.2 Domain-specific application of AI. Educators in Myanmar demonstrate differential application of GenAI across educational domains. The highest mean score was observed for “AI in Education Theories Practices” ($M = 3.53$, $SD = 0.74$), suggesting that educators primarily ground their GenAI use in theoretical frameworks. Comparatively, “AI in Teaching and Learning” showed moderate implementation ($M = 2.66$, $SD = 0.97$), while “AI in Assessment” exhibited the lowest adoption ($M = 2.39$, $SD = 1.04$) (see Table 2).

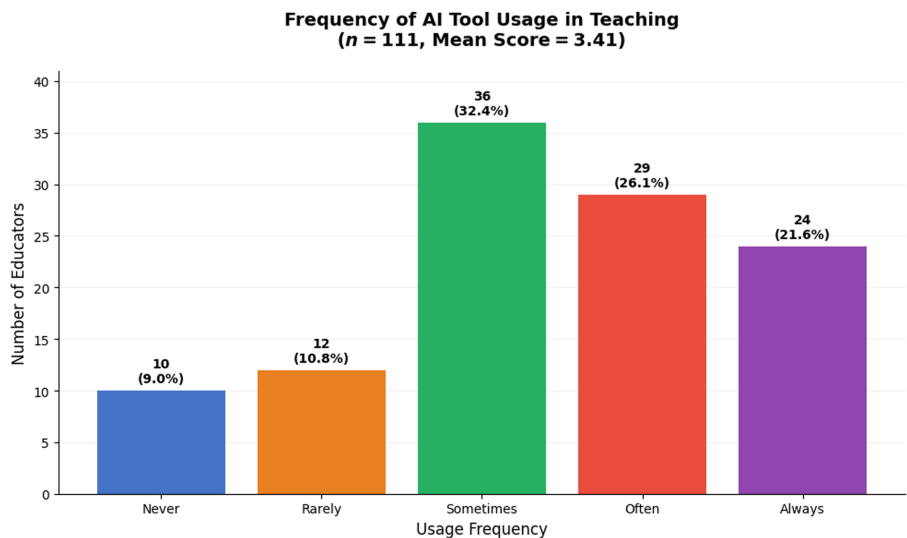


Figure 1. Frequency of AI tool usage in teaching

Table 2. AI use in specific domains

Domain	<i>N</i>	Minimum	Maximum	Mean	Std. Deviation
AI in teaching and learning	111	1.00	5.00	2.66	0.97
AI in assessment	111	1.00	5.00	2.39	1.04
AI in education theories, practices	111	1.00	5.00	3.53	0.74

Note(s): Data collected from 111 educators in Myanmar (*N* = 111). Scale ranges from 1 (lowest usage/implementation) to 5 (highest usage/implementation)

4.1.3 Influence of teaching role on AI integration. The data indicate that dual-role educators (those identifying as both teachers and educators) demonstrate greater GenAI integration across all domains than those identifying solely as teachers. Dual-role educators showed higher mean scores in theoretical applications ($M = 2.88$ vs $M = 2.62$), assessment practices ($M = 2.49$ vs $M = 2.37$) and general educational applications ($M = 3.63$ vs $M = 3.52$) (see Table 3). This finding suggests that educators with broader pedagogical responsibilities may have greater exposure to or interest in implementing GenAI tools across their practice.

4.1.4 Sectoral variations in AI implementation. Significant variations exist in GenAI implementation across different educational sectors in Myanmar. Public education institutions demonstrate the highest levels of GenAI usage in teaching and learning ($M = 2.95$) and assessment practices ($M = 2.79$), while freelance educators show the strongest theoretical grounding ($M = 3.65$) despite having the lowest practical implementation scores (see Table 4). International schools display the most conservative approach to GenAI adoption across all domains, particularly in theoretical applications ($M = 3.05$). These sectoral differences likely reflect variations in institutional resources, policies and professional development opportunities across Myanmar’s educational landscape.

To examine whether the observed sectoral differences were statistically significant, Mann–Whitney *U* tests were conducted comparing the two largest institutional groups – public ($n = 28$) and private ($n = 75$) educators – as the freelance and international school subgroups were too small ($n = 4$ each) to support reliable inferential analysis (for details, see Table 5).

Table 3. Standard deviation showing influence of teaching role on AI integration

Teaching role	Statistic	Teaching and learning theory	Assessment practices	AI in education
Teacher	Mean	2.62	2.37	3.52
	<i>N</i>	94	94	94
	Std. Deviation	0.96	1.06	0.73
Teacher and educator	Mean	2.88	2.49	3.63
	<i>N</i>	17	17	17
	Std. Deviation	1.06	0.95	0.81
Total	Mean	2.66	2.39	3.53
	<i>N</i>	111	111	111
	Std. Deviation	0.97	1.04	0.74

Note(s): Data collected from 111 educators in Myanmar. Scale ranges from 1 (lowest implementation) to 5 (highest implementation)

Table 4. Standard deviation showing sectoral variation in AI integration

Sector	Statistic	Teaching and learning theory	Assessment practices	AI in education
Freelance	Mean	1.73	1.18	3.65
	<i>N</i>	4	4	4
	Std. Deviation	0.88	0.24	0.5
International school	Mean	2.6	1.95	3.05
	<i>N</i>	4	4	4
	Std. Deviation	0.78	0.57	0.62
Private education	Mean	2.61	2.33	3.54
	<i>N</i>	75	75	75
	Std. Deviation	0.93	0.98	0.72
Public education	Mean	2.95	2.79	3.58
	<i>N</i>	28	28	28
	Std. Deviation	1.05	1.15	0.83
Total	Mean	2.66	2.39	3.53
	<i>N</i>	111	111	111
	Std. Deviation	0.97	1.04	0.74

Table 5. Mann–Whitney *U* test results for composite domain scores by sector

Domain	Public <i>M</i> (SD)	Private <i>M</i> (SD)	<i>U</i>	<i>p</i>	<i>r</i>
Teaching and learning	2.95 (1.04)	2.61 (0.93)	1270.0	0.104	0.161
Assessment practices	2.79 (1.14)	2.32 (0.99)	1303.5	0.061	0.185
AI in education (Ethics)	3.58 (0.83)	3.54 (0.72)	1177.5	0.346	0.093
Total	3.12 (0.88)	2.85 (0.72)	1303.0	0.061	0.185

Note(s): *N* = 103 (Public *n* = 28, Private *n* = 75). Effect size $r = Z/\sqrt{N}$

Mann–Whitney *U* tests for the five ethical sub-domains (Beneficence, Non-maleficence, Autonomy, Justice and Explicability) similarly revealed no statistically significant differences between public and private educators (all $p > 0.05$). This pattern indicates that ethical considerations regarding AI use are remarkably consistent across institutional sectors, suggesting that shared cultural and contextual factors – rather than institutional type – shape Myanmar educators’ ethical reasoning about AI.

4.1.5 Influence of teaching experience. The data suggest a nonlinear relationship between teaching experience and GenAI integration. Educators with 3–5 years of experience demonstrate the highest overall implementation ($M = 3.01$), followed by those with 6–10 years ($M = 2.92$), and those with more than 10 years of experience ($M = 2.83$). Novice educators (0–2 years) show the lowest integration levels ($M = 2.56$). This pattern may reflect that mid-career educators have sufficient experience to implement new technologies while remaining receptive to innovation, whereas very new educators may lack the pedagogical foundation to effectively integrate GenAI, and the most experienced educators may demonstrate greater resistance to technological change.

4.1.6 Inferential analysis: experience-level comparisons. Kruskal–Wallis H tests were conducted to examine whether the descriptive patterns across experience levels were statistically significant. As shown in Table 6, no significant differences were found for any composite domain score, though AI in Education (Ethics) approached significance, $H(3) = 7.169$, $p = 0.067$, $\eta^2 = 0.039$.

However, analysis of the five ethical sub-domains revealed one statistically significant finding. Autonomy scores differed significantly across experience levels, $H(3) = 12.846$, $p = 0.005$, $\eta^2 = 0.092$, representing a small-to-medium effect size. As shown in Table 7, endorsement of autonomy-related ethical considerations increased progressively with teaching experience, from $M = 3.00$ ($SD = 0.75$) among novice educators to $M = 4.08$ ($SD = 0.62$) among the most experienced group.

Dunn's post-hoc pairwise comparisons with Bonferroni adjustment indicated that educators with more than 10 years of experience scored significantly higher on Autonomy than those with 0–2 years of experience ($p = 0.003$). No other pairwise comparisons reached significance. Beneficence approached significance ($p = 0.056$), with novice educators scoring notably lower than all other groups, though this did not reach the conventional threshold (see Table 7).

These findings suggest that while most ethical considerations and practical AI implementation patterns are consistent across experience levels, concern for preserving human autonomy in AI-assisted education grows significantly with professional maturity.

Table 6. Kruskal–Wallis H test results for composite domain scores by experience level

Domain	0–2 yrs <i>M</i> (<i>SD</i>)	3–5 yrs <i>M</i> (<i>SD</i>)	6–10 yrs <i>M</i> (<i>SD</i>)	10+ yrs <i>M</i> (<i>SD</i>)	<i>H</i> (3)	<i>p</i>	η^2
Teaching and learning	2.45 (1.08)	2.88 (0.97)	2.64 (0.95)	2.47 (0.94)	3.131	0.372	0.001
Assessment practices	2.21 (1.18)	2.45 (1.03)	2.43 (1.05)	2.30 (1.04)	0.989	0.804	–0.019
AI in education (Ethics)	2.97 (0.83)	3.60 (0.64)	3.57 (0.82)	3.66 (0.51)	7.169	0.067	0.039
Total	2.56 (0.91)	3.01 (0.76)	2.90 (0.78)	2.83 (0.67)	2.992	0.393	0.000

Note(s): $N = 111$ (0–2 years $n = 12$, 3–5 years $n = 34$, 6–10 years $n = 43$, more than 10 years $n = 22$)

Table 7. Kruskal–Wallis H test results for ethical sub-domains by experience level

Ethical principle	0–2 yrs <i>M</i> (<i>SD</i>)	3–5 yrs <i>M</i> (<i>SD</i>)	6–10 yrs <i>M</i> (<i>SD</i>)	10+ yrs <i>M</i> (<i>SD</i>)	<i>H</i> (3)	<i>p</i>	η^2
Beneficence	2.94 (0.99)	3.74 (0.88)	3.79 (0.73)	3.78 (0.69)	7.564	0.056	0.043
Non-maleficence	2.81 (1.27)	3.43 (0.88)	3.37 (1.14)	3.36 (0.76)	3.632	0.304	0.006
Autonomy	3.00 (0.75)	3.64 (0.83)	3.74 (0.87)	4.08 (0.62)	12.846	0.005	0.092
Justice	3.08 (1.16)	3.63 (0.75)	3.42 (0.99)	3.43 (0.82)	2.627	0.453	–0.003
Explicability	3.00 (0.94)	3.57 (0.81)	3.54 (1.06)	3.66 (0.60)	3.758	0.289	0.007

Note(s): $N = 111$. Italic indicates statistical significance at $p < 0.05$

4.1.7 *Learning theory foundations for GenAI implementation.* 4.1.7.1 Behaviorist approach. 69.6% of educators implement GenAI using behaviorist principles at lower intensity levels (1–3), while 30.4% report higher intensity implementation (4–5). The relatively higher adoption of behaviorist principles compared to other approaches may reflect the compatibility between GenAI tools and behaviorist strategies like immediate feedback, reinforcement and structured learning sequences.

4.1.7.2 Cognitive approach. 63.3% of educators apply cognitive learning principles at lower intensity levels, while 36.7% report higher intensity implementation – the highest proportion of high-intensity implementation among all approaches. This suggests that educators in Myanmar find particular value in using GenAI to support cognitive processes such as memory, problem-solving and knowledge organization.

4.1.7.3 Constructivist approach. 72.5% of educators implement constructivist principles at lower intensity levels, with only 27.4% at higher intensity levels. This pattern may reflect challenges in aligning GenAI tools with constructivist ideals of student-directed learning and knowledge construction.

4.1.7.4 Social learning approach. 73.6% of educators utilize social learning theory at lower intensity levels, with just 26.3% at higher levels. This indicates potential difficulties in leveraging GenAI to facilitate meaningful social learning experiences.

4.1.7.5 Active learning approach. 75.6% of educators implement active learning strategies at lower intensity levels, with only 24.3% at higher levels. This suggests limitations in current GenAI applications for supporting student-centered, participatory learning experiences in Myanmar’s educational context.

4.1.7.6 Sociocultural approach. 75.2% of educators apply sociocultural learning principles at lower intensity levels, with just 24.8% at higher levels – the lowest high-intensity implementation among all approaches. This may reflect challenges in using GenAI to address the cultural and contextual dimensions of learning emphasized in sociocultural theory.

4.1.7.7 Overall implementation pattern. Across all learning approaches, 71.6% of implementation occurs at lower intensity levels (1–3), while only 28.4% occurs at higher levels (4–5) (see [Table 8](#)). This distribution suggests a cautious, exploratory approach to GenAI implementation among Myanmar educators, with stronger theoretical grounding than practical integration.

4.2 *Educators’ ethical considerations toward using GenAI*

4.2.1 *Descriptive statistical findings.* Based on responses from 111 Myanmar educators regarding their ethical considerations when using AI in education, measured against a unified ethical framework, the following patterns emerged:

4.2.1.1 Beneficence. Beneficence emerged as the most strongly endorsed ethical principle among participants ($M = 3.68$, $SD = 1.12$). The distribution shows that 58.6% of participants ($n = 65$) agreed or strongly agreed with beneficence considerations, with 28.8% ($n = 32$) strongly agreeing. Only 23.4% ($n = 26$) disagreed or strongly disagreed, while 41.4% ($n = 46$)

Table 8. Total percentages of GenAI use in learning theory foundations

Learning approach	1 (%)	2 (%)	3 (%)	4 (%)	5 (%)	Total (%)
Behaviorist	24.8	18.9	25.9	20	10.4	100
Cognitive	18.9	19.6	24.8	23.4	13.3	100
Constructivist	25	25	22.5	17.3	10.1	100
Social learning	28.8	20.7	24.1	15.3	11	100
Active learning	28.8	23.4	23.4	16	8.3	100
Sociocultural	28.4	25.2	21.6	15.8	9	100
Overall average	25.8	22.1	23.7	18	10.4	100

remained neutral. This indicates that the majority of educators recognize the importance of using AI to benefit their students and enhance educational outcomes.

4.2.1.2 Autonomy. Autonomy received the second highest endorsement ($M = 3.70$, $SD = 1.10$), with 44.1% of participants ($n = 49$) strongly agreeing and an additional 26.1% ($n = 29$) agreeing with autonomy-related considerations. Combined, 70.2% of participants endorsed the importance of maintaining human agency and decision-making in AI-enhanced education. Only 25.2% ($n = 28$) disagreed or strongly disagreed, suggesting strong recognition among educators that AI should support rather than replace human judgment and student autonomy.

4.2.1.3 Explicability. Explicability showed moderate endorsement ($M = 3.52$, $SD = 1.17$), with 62.2% of participants ($n = 69$) agreeing or strongly agreeing that transparency and explainability in AI systems are important. However, 26.1% ($n = 29$) remained neutral, indicating some uncertainty about the importance of understanding how AI systems make decisions. Only 13.5% ($n = 15$) disagreed with explicability considerations.

4.2.1.4 Justice. Justice received moderate support ($M = 3.45$, $SD = 1.12$), with 46.8% of participants ($n = 52$) agreeing or strongly agreeing with justice-related considerations. Notably, 29.7% ($n = 33$) remained neutral, suggesting uncertainty about fairness and equity issues in AI implementation. Approximately 23.4% ($n = 26$) disagreed or strongly disagreed, indicating some educators may not perceive justice as a primary concern in AI use.

4.2.1.5 Non-maleficence. Non-maleficence received the lowest endorsement ($M = 3.33$, $SD = 1.29$), though still above the neutral midpoint. Only 36.9% of participants ($n = 41$) agreed or strongly agreed with non-maleficence considerations, while 31.5% ($n = 35$) disagreed or strongly disagreed. A substantial 27.9% ($n = 31$) remained neutral. This suggests that while some educators recognize the importance of avoiding harm, there may be less awareness or concern about the potential negative consequences of AI implementation in education.

4.2.1.6 Overall pattern. The data reveal that educators most strongly endorse autonomy and beneficence principles, while showing less consensus on non-maleficence, justice and explicability. The relatively high standard deviations across all principles (ranging from 1.10 to 1.29) indicate considerable variability in educators' perspectives on AI ethics. This suggests the need for more comprehensive professional development and discussion around these issues.

4.2.2 Distribution of ethical considerations and acknowledgement of limited knowledge. Based on the Unified Framework of Five Principles for AI in Society (Floridi and Cowsli, 2019), the frequency of ethical considerations mentioned by the 77 respondents was as follows (see Table 9). Participants most frequently expressed concerns about maintaining student independence and avoiding over-reliance on AI. The second most frequent category focused on preventing harm, particularly regarding data privacy and information accuracy. Participants mentioned using AI to enhance rather than replace educational processes. A limited number of participants addressed equity and fairness concerns. Few participants addressed transparency

Table 9. Percentages of distribution of ethical considerations

Ethical consideration	Frequency of mentions	Percentage of 77 respondents
Autonomy	43	55.80
Non-maleficence	20	26.00
Beneficence	16	20.80
Explicability	9	11.70
Justice	8	10.40
Total respondents	77	100
No response	34	30.6% of the total 111 participants

and understanding of AI processes. Several participants explicitly acknowledged their limited understanding of ethical AI use.

Overall, as highlighted in the above data, the ethical considerations that most Myanmar educators consider are mostly aligned with the unified framework of ethical AI.

4.3 Ethical considerations recommended by Myanmar educators

According to the analysis of participants' responses to the open-ended question at the end of the survey (77 responses out of 111), and the audio-recorded workshop discussions from six participants, their responses closely align with the analysis of the statistical trends and reveals the reasoning behind the educators' ethical considerations.

4.3.1 Pedagogical rationale and ethical prioritization. 4.3.1.1 *Beneficence.* The educators clearly recognize AI's potential to promote positive outcomes through efficiency and pedagogical support. Participant A noted, *"I think using AI in teaching, learning, and assessment is a great way to enhance our progress, as it helps reduce the time and energy spent brainstorming how to cover a lesson perfectly."* This response demonstrates an understanding of beneficence through efficiency. Similarly, Participant B highlighted how *"Teachers can learn about new teaching methods and games that children will be interested in by using AI."* These responses show educators appreciate AI's capacity to benefit teaching practice and student engagement. This aligns with Floridi's emphasis on maximizing positive outcomes.

4.3.1.2 *Non-maleficence.* This principle received the strongest attention among educators (20 out of 77 responses), particularly on data privacy and security concerns. Participant A emphasized, *"When using AI in teaching, teachers need to make sure students' data is kept private and secure. They should also be aware of any biases in AI tools, making sure everything is fair and transparent for all students."*

Multiple responses also reveal a concern for data protection. For instance, Participant C warned, *"Students should be particularly careful not to include personal information of students. This is because no one can guarantee security."* Furthermore, a workshop participant also noted potential bias issues, stating that they have *"learned that they often have a bias, especially if they're coming from a Western background."*

These responses demonstrate understanding of non-maleficence, extending beyond immediate harm to long-term privacy risks.

4.3.1.3 *Autonomy.* The highest number of responses (43 out of 77) addressed autonomy, revealing deep concern about maintaining human agency in education. Participant E articulated this well: *"There are many factors to consider, but the most important one is human oversight. Teachers have a deep understanding of their students and the learning process, so they should make the final decisions—such as selecting activities—rather than relying entirely on AI."*

The concern extends to student autonomy as well. Participant G suggested, *"If you are a child, you should explore your own ideas and consult with your friends and then check with AI whether your ideas are right or wrong. So, it would be better if you use AI as a classmate or a roommate."* This metaphor effectively captures the appropriate relationship between students and AI tools.

Perhaps most striking is Participant H's warning that *"if we cannot use this AI effectively in our education system, our teachers' role will be lost and students will become slaves of AI instead of building their own knowledge."* This reflects a deep understanding of autonomy's importance in preserving human agency and intellectual development.

4.3.1.4 *Justice.* While fewer educators explicitly mentioned justice (8 responses), those who did showed clear understanding of equity concerns. Participant A highlighted access inequality, *"For me, I think equal access to every learner and academic integrity will be the most important for conducting AI in the classrooms. So, for example, if one student has the GPT access, but the other student can't access due to Internet connections or due to the cost of the GPT, so it will not be fair to conduct or to use in one classroom."*

Participant C addressed algorithmic bias, noting that “AI systems can reflect biases in their training data. Teachers should critically evaluate AI tools to ensure they do not disadvantage any group of students and promote inclusivity.” This demonstrates awareness of how AI can perpetuate or amplify existing inequalities.

4.3.2 *Explicability, non-disclosure and ethical literacy gaps.* 4.3.2.1 *Explicability.* Regarding the principle of explicability, nine educators emphasized the need for transparency and understanding of how Gen AI works or generates content. For example, Participant A stressed its importance, “As teachers, we must have advanced knowledge of an AI tool before introducing it to the students. Communicating informatively is really important while encouraging the students to use AI as most of the students might wrongly use the tool.”

Similarly, Participant C highlighted the importance of consent and communication in using AI, “Teachers should inform students and parents about how AI is being used in the classroom, the data it collects, and its purpose. Obtaining consent when necessary is critical to building trust.”

4.3.2.2 *Emerging cultural considerations.* Interestingly, the data reveal a tension, which is not explicitly addressed in Florida’s framework – the cultural conflict between transparency and professional authority. Workshop participants expressed reluctance to disclose AI use in their teaching, with one noting, “I don’t usually tell my learners that I have used AI because I am not sure whether they will think that I have up-to-date knowledge with regards to AI or they will think that I used technology too much and that I don’t use my own effort.”

This cultural dimension suggests that implementing ethical AI use in education requires not only following established ethical principles but also addressing contextual factors that may create barriers to transparent practice. The educators’ responses demonstrate sophisticated ethical consideration while revealing real-world implementation challenges that ethical frameworks must consider.

5. Discussion

5.1 Pedagogical foundations, equity risks and AI implementation patterns

The integration of quantitative and qualitative data in this study provides a deep understanding of not only *how* Myanmar educators implement GenAI, but *why* they prioritize specific ethical and pedagogical considerations. The findings suggest that Myanmar educators predominantly ground their GenAI use in a theoretical understanding of educational frameworks before moving to practical applications. This pattern indicates that the use of AI in education should be guided above all by a clear pedagogical rationale, reinforcing the argument that pedagogy must precede technological implementation (Elstad, 2024). This theoretical foundation is consistent with a constructivist approach, where educators first develop their understanding of AI, investigate its pedagogical applications and tackle pressing concerns prior to putting these teaching strategies into practice (Klopfer *et al.*, 2024). Among the specific pedagogical frameworks, the cognitive approach received relatively higher adoption, aligning with research demonstrating that AI excels at supporting cognitive processes through multiple examples, personalized explanations and immediate feedback (Mollick and Mollick, 2023).

Furthermore, the quantitative finding that 71.6% of AI implementation occurs at lower intensity levels reflects a deliberate strategy of considered, limited experimentation (Klopfer *et al.*, 2024). The qualitative insights explain *why* this cautious approach prevails: Myanmar educators are highly aware of the risks of over-reliance on technology. By intentionally limiting AI intensity, these educators exercise a best practice to avoid the specific drawbacks of students using GenAI as a “crutch” without proper pedagogical scaffolding, which can bypass cognitive effort and diminish learning outcomes (Bastani *et al.*, 2024, p. 1). This caution also extends to assessment practices, where implementation was lowest; educators recognize that GenAI has fundamentally disrupted traditional assessment methods and are deliberately limiting its use to maintain quality evaluation (Hodges and Kirschner, 2024).

The inferential analysis strengthens these observations while adding an important nuance regarding institutions. The typical assumption in global discourse is that resource-constrained schools will have less access to GenAI, widening the digital divide. However, our data reveal that public education institutions in Myanmar demonstrate similar practical implementation like private sector institutions. This suggests that public educators as well as educators from private sectors are innovatively using free GenAI tools to overcome existing resource gaps. Consequently, the new equity risk is one of quality and safety for vulnerable students (Klopfer *et al.*, 2024). Crucially, Mann–Whitney *U* tests revealed no statistically significant differences between public and private sector educators on any composite domain score or ethical sub-domain (all $p > 0.05$). This null finding proves that shared cultural norms, professional values and contextual constraints in the Global South operate across institutional boundaries and drive AI adoption more strongly than sector-specific resources.

Similarly, Kruskal–Wallis *H* tests revealed no significant differences across experience levels for any composite domain score (all $p > 0.05$), reinforcing the conclusion that practical AI implementation patterns are broadly consistent regardless of professional experience. The one exception – the significant growth of autonomy concerns with experience – is discussed in Section 5.2.

5.2 Ethical reasoning, cultural wisdom and the literacy gap

The findings indicate that Myanmar educators demonstrate a strong intuitive grasp of the ethical principles commonly discussed in global AI ethics frameworks, even without formal exposure to the unified frameworks established by Floridi *et al.* (2018) and Floridi and COWls (2019).

The quantitative data confirm this pattern: beneficence ($M = 3.68$) and autonomy ($M = 3.70$) received the strongest endorsement on the Likert scale, while non-maleficence received the lowest ($M = 3.33$). This ranking mirrors the qualitative frequency data, where autonomy dominated open-ended responses (43 of 77) and justice and explicability were least cited. The convergence of quantitative endorsement levels and qualitative response frequencies strengthens the validity of the finding that Myanmar educators' ethical priorities cluster around principles most directly connected to their daily pedagogical practice.

When discussing beneficence, participants framed AI's role as supportive rather than transformative, reflecting a culturally grounded wisdom that technology should strengthen, rather than erode, the teacher's central pedagogical role (Klopfer *et al.*, 2024). Similarly, the prevalence of non-maleficence concerns demonstrates that Myanmar educators share universal ethical priorities regarding data privacy and security (Funa and Gabay, 2025). However, their skepticism about claims of complete security perfectly mirrors the principle of "capability caution" (Floridi and COWls, 2019, p. 6), arising organically from their lived experiences in fragile technological environments rather than from formal ethical training.

The inferential statistical analyses provide crucial depth to these qualitative patterns. While most ethical considerations showed no statistically significant variation across demographics, the Kruskal–Wallis *H* tests revealed that autonomy was the only ethical principle that differed significantly across experience levels, $H(3) = 12.846$, $p = 0.005$, $\eta^2 = 0.092$. Endorsement of autonomy increased progressively from novice educators ($M = 3.00$) to those with more than 10 years of experience ($M = 4.08$). The qualitative data explains exactly *why*: highly experienced educators possess a nuanced fear that without proper human oversight, teachers' roles will be lost and students will become "slaves of AI". This aligns with the concept of "meta-autonomy" – the capacity to decide which decisions to delegate to AI – which is cultivated through professional practice (Floridi *et al.*, 2018, p. 698).

That autonomy alone reaches statistical significance, while principles like justice and explicability show no variation by experience, confirms a critical vulnerability. The substantial non-response rate regarding justice and explicability points to a systemic gap in formal ethical

AI literacy rather than a deliberate deprioritization. This ethical AI literacy deficit highlights the heavy cost of excluding the Global South from the global AI discourse, underscoring the urgent need to co-create localized ethical guidelines.

5.3 Cultural frictions: global transparency vs local authority

The most significant finding of this study highlights an epistemic and cultural clash between imported AI ethics and local realities. While dominant AI ethics frameworks developed in the Global North position explicit transparency, acknowledgment and explicability as non-negotiable considerations for AI deployment (Adams, 2021; Floridi and Cows, 2019; Moorhouse *et al.*, 2023), participating Myanmar educators frequently choose *not* to inform students when they use AI to create lesson content.

This decision is deeply tied to Myanmar's cultural context. In Myanmar, the tradition of considering teachers as ultimate knowledge guardians – reinforced by hierarchical homage-paying ceremonies like the *gadaw* – maintains a highly teacher-centered classroom culture (Lall and South, 2014; Lwin, 2000; Seekins, 2017). Within this dynamic, where students are expected to show deep respect and view educators as the primary source of knowledge, admitting reliance on an external AI tool is perceived as a direct threat to a teacher's professional status and intellectual authority (Lall and South, 2014; Lwin, 2000). Because existing global frameworks are typically imported from the Global North, they do not account for these specific sociocultural dynamics (Adams, 2021; Vijayakumar, 2024).

By demanding transparency without understanding local authority structures, universal frameworks inadvertently force Myanmar educators into a reactive posture where they must secretly navigate AI integration rather than openly model its responsible use. When teachers are not actively involved in shaping the policies that govern their classrooms, their professional judgment and autonomy are undermined, limiting opportunities for meaningful, transparent change (Kennedy and Castek, 2025). Addressing this epistemic friction requires moving away from the top-down imposition of external standards and toward the development of inclusive, context-specific ethical frameworks that respect the unique educational realities of the Global South (Vijayakumar, 2024).

In summary, the mixed-methods analysis provides a comprehensive picture of AI application in Myanmar's educational system. The quantitative data established that while practical AI implementation is relatively consistent across institutional sectors, an educator's concern for human autonomy deepens significantly with professional experience. The qualitative data deepened these results by uncovering the "why" behind the phenomena: educators intuitively prioritize autonomy to protect their students' intellectual agency and avoid the use of AI as a pedagogical crutch, yet their practical implementation is hindered by a systemic deficit in formal ethical AI literacy. Most importantly, the research identified a fundamental conflict between global AI principles and local culture. The cultural imperative to maintain teacher authority and respect in Myanmar directly contradicts Western mandates for transparency in AI use, demonstrating that implementing ethical AI use in education requires addressing contextual barriers that prevent transparent practice.

6. Conclusion

This research presents the perspectives of Myanmar educators on the ethics of GenAI in education, demonstrating that their ethical priorities reflect both deeply rooted cultural values and practical pedagogical needs. While these educators place a strong emphasis on autonomy and beneficence, which support the teacher's central role in the classroom and bring tangible benefits to everyday practice, they are significantly less engaged with ethical principles like justice and explicability. As the findings establish, this does not indicate a deliberate deprioritization, but rather stems from a systemic ethical AI literacy gap exacerbated by the

Global South's exclusion from the main international discussions where such concepts are defined (Adams, 2021; Nemorin, 2024).

Furthermore, the study highlights a cultural tension: while Western ethical models mandate strict transparency in AI use, many educators in Myanmar choose to hide their AI use to preserve their intellectual authority, viewing unquestioned teacher expertise as essential to their professional identity and established classroom hierarchies (Lall and South, 2014). The unadapted imposition of universal, Global North ethical rules to such diverse settings risks ignoring important local realities and forcing educators into a reactive, rather than proactive, posture toward AI adoption.

6.1 Implications and recommendations for practice

To address these issues, ethical guidelines for AI in education cannot simply be imported. They must be co-developed with the educators who understand the socioeconomic and cultural realities of their own contexts. To provide practical guidance on operationalizing these goals, we recommend the following steps based on our findings:

- *Contextualized policy co-design:* To address the cultural friction surrounding transparency, schools and districts should establish formal channels for teachers to actively participate in AI policy discussions (Kennedy and Castek, 2025). Educational institutions should host localized co-design workshops where educators draft their own acceptable-use guidelines. This will ensure that policies respect local authority structures and classroom management needs, rather than blindly adopting outsider mandates for disclosure.
- *Targeted professional development:* To address the critical ethical AI literacy gap, professional development should move beyond basic digital skills to explicitly cover algorithmic justice, bias and data privacy. Because our quantitative data show that autonomy concerns grow significantly with professional maturity, schools should establish structured mentorship programs. Highly experienced teachers should guide novice educators in developing “meta-autonomy” – the critical pedagogical judgment required to safely decide when and how to delegate tasks to AI tools (Floridi *et al.*, 2018, p. 698).
- *Inclusive global representation:* To counter the marginalization of Global South perspectives, educators should be included in international debates to ensure ethical standards reflect genuine diversity. Regional educational bodies and civil society organizations should form consortiums to actively engage with international policymakers (e.g. UNESCO). These consortiums must demand that global ethical frameworks incorporate pluriversal views, ensuring guidelines are adaptable to local realities rather than reinforcing a single, Eurocentric perspective (Nemorin, 2024; Vijayakumar, 2024).

Further research should compare different regions across the Global South and examine longitudinal changes, as teachers' understanding of AI ethics will likely evolve with new experiences and targeted training. Finally, we argue that centering these marginalized voices in global debates is an absolute necessity for creating AI governance frameworks that are both equitable and effective.

Author contribution

All authors contributed equally and share first authorship. The authors are listed alphabetically by surname.

References

- Adams, R. (2021), "Can artificial intelligence be decolonized?", *Interdisciplinary Science Reviews*, Vol. 46 Nos 1-2, pp. 176-197, doi: [10.1080/03080188.2020.1840225](https://doi.org/10.1080/03080188.2020.1840225).
- Akila, D., Garg, H., Pal, S. and Jeyalakshmi, S. (2023), "Research on recognition of students attention in offline classroom-based on deep learning", *Education and Information Technologies*, Vol. 29 No. 6, pp. 6865-6893, doi: [10.1007/s10639-023-12089-6](https://doi.org/10.1007/s10639-023-12089-6).
- Anderson, J. and Mahapatra, S. (2024), "From 'difficult circumstances' to the 'Global South': an appraisal of terms used to describe disadvantage in education in the Global South", *Fortell*, Vol. 49, pp. 7-23.
- Ausubel, D.P. (1968), *Educational Psychology: A Cognitive View*, Holt, Rinehart & Winston, New York.
- Bandura, A. (1977), *Social Learning Theory*, Englewood Cliffs, NJ: Prentice-Hall, available at: https://www.ascib.ase.ro/mps/Bandura_SocialLearningTheory.pdf
- Bastani, H., Bastani, O., Sungu, A., Ge, H., Kabakci, Ö. and Mariman, R. (2024), "Generative AI can harm learning", *SSRN Electronic Journal*, Vol. 112 No. 26, doi: [10.2139/ssrn.4895486](https://doi.org/10.2139/ssrn.4895486).
- Beauchamp, T.L. and Childress, J.F. (2013), *Principles of Biomedical Ethics*, 7th ed., Oxford University Press, Oxford.
- Borenstein, J. and Howard, A. (2021), "Emerging challenges in AI and the need for AI ethics education", *AI Ethics*, Vol. 1, pp. 61-65, doi: [10.1007/s43681-020-00002-7](https://doi.org/10.1007/s43681-020-00002-7).
- Bruner, J.S. (1966), *Toward a Theory of Instruction*, Harvard University Press, Cambridge.
- Chen, B., Zhu, X. and Díaz del Castillo, H.F. (2023), "Integrating generative AI in knowledge building", *Computers and Education: Artificial Intelligence*, Vol. 5, 100184, doi: [10.1016/j.caeai.2023.100184](https://doi.org/10.1016/j.caeai.2023.100184).
- Creswell, J.W. (2014), *Research Design: Qualitative, Quantitative, and Mixed Methods Approaches*, 4th ed., Sage Publications, Thousand Oaks.
- Cumbo, B. and Selwyn, N. (2021), "Using participatory design approaches in educational research", *International Journal of Research and Method in Education*, Vol. 45 No. 1, pp. 60-72, doi: [10.1080/1743727X.2021.1902981](https://doi.org/10.1080/1743727X.2021.1902981).
- Dewey, J. (1997), *Experience and Education*, Simon & Schuster, New York.
- Elstad, E. (2024), "AI in education: rationale, principles, and instructional implications", doi: [10.48550/arxiv.2412.12116](https://doi.org/10.48550/arxiv.2412.12116).
- European Group on Ethics in Science and New Technologies (2018), *Statement on Artificial Intelligence, Robotics and 'autonomous' Systems*, European Commission, doi: [10.2777/531856](https://doi.org/10.2777/531856).
- Fereday, J. and Muir-Cochrane, E. (2006), "Demonstrating rigor using thematic analysis: a hybrid approach of inductive and deductive coding and theme development", *International Journal of Qualitative Methods*, Vol. 5 No. 1, pp. 80-92, doi: [10.1177/160940690600500107](https://doi.org/10.1177/160940690600500107).
- Field, A. (2018), *Discovering Statistics Using IBM SPSS Statistics*, 5th ed., Sage Publications, Los Angeles.
- Floridi, L. (2013), *The Ethics of Information*, Oxford University Press, Oxford.
- Floridi, L. (2023), *The Ethics of Artificial Intelligence: Principles, Challenges, and Opportunities*, Oxford University Press, Oxford.
- Floridi, L. and Cowls, J. (2019), "A unified framework of five principles for AI in society", *Harvard Data Science Review*, Vol. 1 No. 1, pp. 1-15, doi: [10.1162/99608f92.8cd550d1](https://doi.org/10.1162/99608f92.8cd550d1).
- Floridi, L., Cowls, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Luetge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P. and Vayena, E. (2018), "AI4People—an ethical framework for a good AI society: opportunities, risks, principles, and recommendations", *Minds and Machines*, Vol. 28 No. 4, pp. 689-707, doi: [10.1007/s11023-018-9482-5](https://doi.org/10.1007/s11023-018-9482-5).

- Funa, A.A. and Gabay, R.A.E. (2025), "Policy guidelines and recommendations on AI use in teaching and learning: a meta-synthesis study", *Social Sciences and Humanities Open*, Vol. 11, 101221, doi: [10.1016/j.ssaho.2024.101221](https://doi.org/10.1016/j.ssaho.2024.101221).
- Future of Life Institute (2017), "Asilomar AI principles", available at: <https://futureoflife.org/ai-principles/>
- Hodges, C.B. and Kirschner, P.A. (2024), "Innovation of instructional design and assessment in the age of generative artificial intelligence", *TechTrends*, Vol. 68 No. 1, pp. 195-199, doi: [10.1007/s11528-023-00926-x](https://doi.org/10.1007/s11528-023-00926-x).
- House of Lords Artificial Intelligence Committee (2018), *AI in the UK: Ready, Willing and Able?*, House of Lords. Publications, available at: parliament.uk/pa/ld201719/ldselect/ldai/100/100.pdf (accessed 1 June 2025).
- Johnson, R.B. and Onwuegbuzie, A.J. (2004), "Mixed methods research: a research paradigm whose time has come", *Educational Researcher*, Vol. 33 No. 7, pp. 14-26, doi: [10.3102/0013189X033007014](https://doi.org/10.3102/0013189X033007014).
- Kasneci, E., Sessler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., Gasser, U., Groh, G., Günnemann, S., Hüllermeier, E., Krusche, S., Kutyniok, G., Michaeli, T., Nerdel, C., Pfeffer, J., Poquet, O., Sailer, M., Schmidt, A., Seidel, T., Stadler, M., Weller, J., Kuhn, J. and Kasneci, G. (2023), "ChatGPT for good? On opportunities and challenges of large language models for education", *Learning and Individual Differences*, Vol. 103, 102274, doi: [10.1016/j.lindif.2023.102274](https://doi.org/10.1016/j.lindif.2023.102274).
- Kennedy, K.J. and Castek, J.M. (2025), "Empowering ELA teachers: recommendations for teacher education in the AI Era", *Contemporary Issues in Technology and Teacher Education*, Vol. 25 No. 1, p. n1, available at: <https://citejournal.org/volume-25/issue-1-25/english-language-arts/empowering-ela-teachers-recommendations-for-teacher-education-in-the-ai-era/> (accessed 1 June 2025).
- Klopfer, E., Reich, J., Abelson, H. and Breazeal, C. (2024), "Generative AI and K-12 education: an MIT perspective", *An MIT Exploration of Generative AI*. doi: [10.21428/e4baedd9.81164b06](https://doi.org/10.21428/e4baedd9.81164b06).
- Koh, J., Cowling, M., Jha, M. and Sim, K.N. (2023), "The human teacher, the AI teacher and the AId-teacher relationship", *Journal of Higher Education Theory and Practice*, Vol. 23 No. 17, doi: [10.33423/jhetp.v23i17.6543](https://doi.org/10.33423/jhetp.v23i17.6543).
- Lall, M. and South, A. (2014), "Comparing models of non-state ethnic education in Myanmar: the Mon and Karen national education regimes", *Journal of Contemporary Asia*, Vol. 44 No. 2, pp. 298-321, doi: [10.1080/00472336.2013.823534](https://doi.org/10.1080/00472336.2013.823534).
- Lave, J. and Wenger, E. (1991), *Situated Learning: Legitimate Peripheral Participation*, Cambridge University Press, Cambridge.
- Lwin, T. (2000), "Education in Burma (1945-2000). Thinking classroom foundation", available at: http://www.thinkingclassroom.org/uploads/4/3/9/0/43900311/1._lwin_t._2000_education_in_burma__1945-2000__2000__english_.pdf (accessed 11 May 2025).
- Mancilla-Caceres, J.F. and Estrada-Villalta, S. (2022), "The ethical considerations of AI in Latin America", *Digital Society*, Vol. 1 No. 16, doi: [10.1007/s44206-022-00018-y](https://doi.org/10.1007/s44206-022-00018-y).
- Mejeh, M., Hagenauer, G. and Gläser-Zikuda, M. (2023), "Mixed methods research on learning and instruction—meeting the challenges of multiple perspectives and levels within a complex field", *Forum for Qualitative Social Research*, Vol. 24 No. 1, doi: [10.17169/fqs-24.1.3989](https://doi.org/10.17169/fqs-24.1.3989).
- Mollick, E.R. and Mollick, L. (2023), "Using AI to implement effective teaching strategies in classrooms: five strategies, including prompts", *SSRN Electronic Journal*. doi: [10.2139/ssrn.4391243](https://doi.org/10.2139/ssrn.4391243).
- Monasterio Astobiza, A., Ausín, T., Liedo, B., Toboso, M., Aparicio, M. and López, D. (2022), "Ethical governance of AI in the Global South: a human rights approach to responsible use of AI", *Proceedings*, Vol. 81 No. 1, doi: [10.3390/proceedings2022081136](https://doi.org/10.3390/proceedings2022081136).
- Moorhouse, B.L., Yeo, M.A. and Wan, Y. (2023), "Generative AI tools and assessment: guidelines of the world's top-ranking universities", *Computers and Education Open*, Vol. 5, 100151, doi: [10.1016/j.cao.2023.100151](https://doi.org/10.1016/j.cao.2023.100151).

- Moss, P.A. and Haertel, E.H. (2016), "Engaging methodological pluralism", in Gitomer, D.H. and Bell, C.A. (Eds), *Handbook of Research on Teaching*, American Educational Research Association, Washington, pp. 127-247.
- Moussa, M.K. (2024), "Towards reliable utilization: an instructional design model for integrating generative pre-trained transformer (GPT) in education", in Al-Marzouqi, A., Salloum, S.A., Al-Saidat, M., Aburayya, A. and Gupta, B. (Eds), *Artificial Intelligence in Education: the Power and Dangers of ChatGPT in the Classroom*, 1st ed., Springer, Cham, Vol. 144, pp. 481-496.
- Ncube, C.N. and Tawanda, T. (2025), "Critical digital pedagogy for contemporary transformative practices in the Global South: a literature review", *Cogent Education*, Vol. 12 No. 1, doi: [10.1080/2331186X.2025.2523133](https://doi.org/10.1080/2331186X.2025.2523133).
- Nemorin, S. (2024), "Towards decolonising the ethics of AI in education", *Globalisation, Societies and Education*, Vol. 24 No. 4, pp. 1-13, doi: [10.1080/14767724.2024.2333821](https://doi.org/10.1080/14767724.2024.2333821).
- Nunnally, J.C. (1978), *Psychometric Theory*, 2nd ed., McGraw-Hill, New York.
- Piaget, J. (2007), *The Psychology of Intelligence*, Routledge, London; NY.
- Ricaurte, P. (2019), "Data epistemologies, the coloniality of power, and resistance", *Television and New Media*, Vol. 20 No. 4, pp. 350-365, doi: [10.1177/1527476419831640](https://doi.org/10.1177/1527476419831640).
- Roche, C., Lewis, D. and Wall, P.J. (2021), "Artificial intelligence ethics: an inclusive global discourse?", *arXiv*. doi: [10.48550/arXiv.2108.09959](https://doi.org/10.48550/arXiv.2108.09959).
- Seekins, D.M. (2017), *Historical Dictionary of Burma (Myanmar)*, Rowman & Littlefield, Lanham, MD.
- Shahriari, K. and Shahriari, M. (2017), "IEEE standard review — ethically aligned design: a vision for prioritizing human wellbeing with artificial intelligence and autonomous systems", *2017 IEEE Canada International Humanitarian Technology Conference (IHTC)*, pp. 197-201, doi: [10.1109/ihc.2017.8058187](https://doi.org/10.1109/ihc.2017.8058187).
- Skinner, B.F. (1938), *The Behavior of Organisms: an Experimental Analysis*, B. F. Skinner Foundation, Cambridge.
- Terwiesch, C. (2023), *Would Chat GPT3 Get a Wharton MBA? A Prediction Based on its Performance in the Operations Management Course*, available at: <https://mackinstitute.wharton.upenn.edu/wp-content/uploads/2023/01/Christian-Terwiesch-Chat-GTP.pdf>.
- Tun, P.W. and Smith, R. (2026), "English teaching online, in politically difficult circumstances", *ELT Journal*, Vol. 80 No. 3, pp. 331-343, doi: [10.1093/elt/ccag019](https://doi.org/10.1093/elt/ccag019).
- UNESCO (2024), "UNESCO leads ethical AI discussions at AI Global South Summit 2024", available at: <https://www.unesco.org/en/articles/unesco-leads-ethical-ai-discussions-ai-global-south-summit-2024> (accessed 1 June 2025).
- Université de Montréal (2017), *Montreal Declaration for Responsible AI*, available at: <https://recherche.umontreal.ca/english/strategic-initiatives/montreal-declaration-for-a-responsible-ai/>
- Valdivieso, T. and González, O. (2025), "Generative AI tools in Salvadoran higher education: balancing equity, ethics, and knowledge management in the Global South", *Education Sciences*, Vol. 15 No. 2, doi: [10.3390/educsci15020214](https://doi.org/10.3390/educsci15020214).
- Vijayakumar, A. (2024), "AI ethics for the Global South: perspectives, practicalities, and India's role", (RIS discussion Paper No. 296). Research and Information System for Developing Countries, available at: <https://www.ris.org.in/sites/default/files/Publication/DP-296-Anupama-Vijayakumar.pdf> (accessed 1 June 2025).
- von Glasersfeld, E. (1995), *Radical Constructivism: A Way of Knowing and Learning*, Falmer Press, available at: <https://files.eric.ed.gov/fulltext/ED381352.pdf>
- Vygotsky, L.S. (1978), *Mind in Society: The Development of Higher Psychological Processes*, Harvard University Press, Cambridge.
- Watson, J.B. (1998), *Behaviorism*, (1st ed.), Routledge, doi: [10.4324/9781351314329](https://doi.org/10.4324/9781351314329)
- Wenger, E. (1998), *Communities of Practice: Learning, Meaning, and Identity*, Cambridge University Press, Cambridge.

Xia, Q., Weng, X., Ouyang, F., Lin, T.J. and Chiu, T.K.F. (2024), "A scoping review on how generative artificial intelligence transforms assessment in higher education", *International Journal of Educational Technology in Higher Education*, Vol. 21 No. 1, 40, doi: [10.1186/s41239-024-00468-z](https://doi.org/10.1186/s41239-024-00468-z).

Zawacki-Richter, O., Marín, V.I., Bond, M. and 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*, Vol. 16 No. 39, pp. 1-27, doi: [10.1186/s41239-019-0171-0](https://doi.org/10.1186/s41239-019-0171-0).

Corresponding author

Thu Thu Naing can be contacted at: thuthunaing26@gmail.com