

Academic Motivation and Integrity in the Era of Generative AI: Critical Reflections on Students' Dependence on AI Tools

Wan Nabilah Syafiqah Wan Abd Rahim^{1*}, & Nurul Aisyah Kamrozzaman², Sofia Elias³

¹Faculty of Education and Humanities, UNITAR International University, Malaysia

²Faculty of Education and Humanities, UNITAR International University, Malaysia

³Faculty of Education and Humanities, UNITAR International University, Malaysia

*Corresponding author:

DOI : <https://doi.org/10.47772/IJRISS.2026.102400005>

Received: 20 August 2026; Accepted: 25 August 2026; Published: 25 September 2026

ABSTRACT

The rapid diffusion of generative artificial intelligence (GenAI) tools such as ChatGPT has fundamentally altered how university students complete academic tasks, raising urgent questions about what it now means to learn, to know, and to act with integrity. This article presents a critical reflection on students' growing dependence on GenAI, focusing on two interrelated concerns: the transformation of academic motivation and the erosion of academic integrity. Drawing on self-regulated learning theory, achievement goal theory, self-determination theory, and the literature on cheating and academic dishonesty, the paper argues that habitual delegation of intellectual work to AI risks shifting students' motivational orientation away from mastery of knowledge toward performance outcomes and, more troublingly, toward an emerging orientation of expedience in which speed and convenience displace both mastery and performance goals. The paper further contends that GenAI destabilises inherited definitions of cheating by creating a wide grey zone between legitimate assistance and dishonest outsourcing, thereby demanding a reconceptualisation of integrity as an internal disposition rather than mere rule compliance. The critical reflection develops five theoretical propositions concerning motivational displacement, metacognitive offloading, the redefinition of dishonesty, the formation of academic character, and institutional responsibility. The discussion elaborates the long-term implications of AI dependence for students' epistemic agency, intellectual virtues, and readiness for lifelong learning, and it aligns these concerns with United Nations Sustainable Development Goal 4 on quality education. The paper concludes that the central challenge of the GenAI era is not detecting misconduct but cultivating learners who choose effortful learning when effortless output is available, and it offers normative directions for pedagogy, assessment design, and institutional policy.

Keywords: generative artificial intelligence; academic integrity; academic motivation; self-regulated learning; Quality Education

INTRODUCTION

The public release of ChatGPT in November 2022, followed by a proliferation of increasingly capable generative artificial intelligence (GenAI) systems, has produced one of the most consequential disruptions to higher education in decades (Kasneci et al., 2023; Rudolph et al., 2023). Within months of its launch, students across the world were using large language models to draft essays, solve problem sets, summarise readings, generate code, and answer examination questions (Chan & Hu, 2023; Ngo, 2023). Unlike earlier educational technologies, which primarily supported learning processes, GenAI is able to produce the final artefacts of learning: coherent, original-seeming text and solutions that are largely indistinguishable from competent student work (Cotton et al., 2024; Nikolic et al., 2023). This capability strikes at the heart of the assessment logic on which universities have long depended, namely the assumption that a submitted artefact is a valid proxy for the knowledge and skill of the person who submitted it.

Much of the early scholarly and institutional response to this disruption has been framed in policing terms: how to detect AI-generated text, how to redesign assessments to be AI-proof, and how to update misconduct regulations (Perkins, 2023; Sullivan et al., 2023). While these are legitimate concerns, this article argues that the detection-and-deterrence framing captures only the surface of the problem. Beneath the visible question of whether students are cheating lies a deeper and more consequential question: what is happening to students' motivation to learn, and to their character as learners, when a substantial portion of their intellectual labour can be delegated to a machine? If a student can produce an acceptable assignment in minutes without engaging with the underlying material, the issue is not only whether the act violates a rule, but what repeated acts of this kind do to the student's orientation toward knowledge, effort, and honesty over time (Crawford et al., 2023; Zhai et al., 2024).

There is a substantial and rapidly growing empirical literature on students' adoption of GenAI, their perceptions of its benefits and risks, and the institutional policies emerging in response (Abbas et al., 2024; Chan, 2023; Farrelly & Baker, 2023; Lo, 2023). What remains comparatively underdeveloped is sustained normative and theoretical reflection that connects this new technological reality to established bodies of theory on academic motivation, self-regulated learning (SRL), and academic dishonesty. Decades of research have shown that motivation and integrity are intimately linked: students who pursue mastery goals cheat less, students who pursue performance and avoidance goals cheat more, and classroom structures that emphasise grades over understanding elevate dishonesty (Anderman & Koenka, 2017; Krou et al., 2021; Murdock & Anderman, 2006). GenAI enters precisely into this motivational field, dramatically lowering the cost of performance without mastery. Understanding its implications therefore requires theoretical integration rather than only empirical description.

This article responds to that need by offering a critical reflection, a genre of scholarship that evaluates a phenomenon normatively and theoretically rather than testing hypotheses empirically. The objective of the paper is threefold. First, it examines the integrity dimension of AI dependence, asking what learning means when large parts of task completion are AI-assisted, and how the boundary between assistance and dishonesty should be redrawn.

Second, it analyses how dependence on AI may reshape students' motivational orientations, arguing that the classic distinction between mastery and performance goals must now be supplemented by an orientation of expedience, a disposition to prioritise what is easy and fast. Third, it reflects on the long-term implications of these shifts for academic character, epistemic agency, and the purposes of higher education, and it derives implications for pedagogy, assessment, and policy. Throughout, the paper draws on the literatures on self-regulated learning, achievement goal theory, self-determination theory, cognitive offloading, and academic dishonesty as its theoretical backdrop.

The significance of this reflection extends beyond individual classrooms. The United Nations Sustainable Development Goal 4 (SDG 4) commits the global community to ensuring inclusive and equitable quality education and promoting lifelong learning opportunities for all. Quality education, in this sense, is not merely the production of credentials; it is the formation of people who can think, judge, and continue learning independently (UNESCO, 2023). If widespread AI dependence hollows out the processes through which such capacities are formed, higher education may continue to certify graduates while quietly failing to educate them, an outcome that would undermine SDG 4 even as enrolment and completion statistics improve. The stakes of the present reflection are therefore both ethical and developmental.

The remainder of the paper is organised as follows. The literature review synthesises scholarship on GenAI in higher education, self-regulated learning, motivational orientations, and academic dishonesty. The methodology section describes the approach of the critical reflection. The paper then presents five theoretical propositions that structure the reflection, followed by an extended discussion that elaborates each proposition in detail, considers counterarguments, and derives implications. The conclusion summarises the contribution, acknowledges limitations, and identifies directions for future research.

LITERATURE REVIEW

Generative AI and the Changing Landscape of Academic Work

GenAI systems, built on large language models trained on vast text corpora, are capable of producing fluent prose, working code, structured arguments, and plausible analyses in response to natural language prompts (Dwivedi et al., 2023; Kasneci et al., 2023). Reviews of the emerging literature identify a consistent set of educational affordances: personalised explanation, instant feedback, language support for non-native speakers, brainstorming assistance, and the democratisation of tutoring-like support (Adiguzel et al., 2023; Lo, 2023; Rasul et al., 2023). Students themselves report largely positive perceptions, valuing the speed, availability, and non-judgemental character of AI assistance (Chan & Hu, 2023; Ngo, 2023). Some scholars have described GenAI as a potential more knowledgeable other in the Vygotskian sense, capable of scaffolding learning within a student's zone of proximal development (Stojanov, 2023).

At the same time, the literature documents serious concerns. Benchmarking studies show that GenAI can pass or perform credibly on a wide range of university assessments, undermining the validity of unsupervised written tasks (Nikolic et al., 2023). Scholars warn of hallucinated content, embedded bias, opacity, and threats to fairness and accountability (Memarian & Doleck, 2023; Michel-Villarreal et al., 2023). Most relevant to the present reflection is the growing body of work on over-reliance. Zhai et al. (2024) synthesised evidence that habitual dependence on AI dialogue systems is associated with diminished decision-making, critical thinking, and analytical reasoning. Abbas et al. (2024) found that academic workload and time pressure drive GenAI use, and that heavier use predicts procrastination, memory loss, and lower academic performance. Survey and experimental work with knowledge workers similarly suggests that greater confidence in AI is associated with reduced self-reported critical thinking effort (Gerlich, 2025; Lee et al., 2025). Together, these findings indicate that the question is no longer whether students will use GenAI, but how patterns of use will shape their cognitive and motivational development (Darvishi et al., 2024).

Self-Regulated Learning and the Value of Effortful Process

Self-regulated learning theory provides the most direct lens for understanding what is at stake when task completion is delegated to AI. SRL describes learning as a cyclical process in which learners set goals, plan strategies, monitor their comprehension and progress, and reflect on outcomes in ways that inform subsequent learning (Panadero, 2017; Zimmerman, 2002). Classic work by Pintrich and De Groot (1990) established that motivational beliefs, particularly self-efficacy and intrinsic value, are tightly interwoven with the use of cognitive and metacognitive strategies. Self-efficacy itself, in Bandura's (1997) account, is built primarily through mastery experiences: authentic episodes of overcoming difficulty through one's own effort. Recent Malaysian evidence reaffirms these classic relationships in contemporary blended learning contexts, showing strong positive associations between self-regulated strategies and intrinsic motivation and self-efficacy among higher education students (Rifin et al., 2025).

The relevance of SRL to the GenAI debate is twofold. First, the theory implies that the pedagogical value of an academic task lies largely in the process it demands, not the artefact it yields. Writing an essay forces retrieval, organisation, self-monitoring, and revision; these are the mechanisms through which understanding consolidates. When AI produces the artefact, the process is skipped, and with it the learning (Bearman & Ajjaw, 2023; Zhai et al., 2024). Second, SRL research shows that regulation itself is a learned capability that develops through practice. Scholars of hybrid human-AI learning warn that if AI systems take over the regulatory functions of planning, monitoring, and evaluating, learners may fail to develop these capabilities at all, a risk that makes the deliberate distribution of control between learner and machine a central design question (Järvelä et al., 2023; Molenaar, 2022). Darvishi et al. (2024), in a large field experiment, found that students who received AI assistance came to rely on it and did not internalise the supported skills once the assistance was withdrawn, illustrating precisely this concern.

This concern connects to the older literature on cognitive offloading, which shows that people readily delegate memory and thought to external tools and subsequently remember and process less themselves

(Risko & Gilbert, 2016; Sparrow et al., 2011). GenAI radically expands the scope of what can be offloaded: no longer only storage and retrieval, but analysis, synthesis, argumentation, and composition. Offloading is not inherently harmful; experts offload routine operations to free capacity for higher-order work. The developmental problem arises when novices offload the very operations they are supposed to be acquiring (Gerlich, 2025; Markauskaite et al., 2022).

Motivational Orientations: Mastery, Performance, and the Rise of Expedience

Achievement goal theory distinguishes between mastery goals, in which success is defined as developing competence and understanding, and performance goals, in which success is defined as demonstrating competence relative to others or securing favourable judgements such as grades (Ames, 1992; Dweck, 1986; Elliot & McGregor, 2001). Decades of research associate mastery orientation with deeper strategy use, greater persistence, higher intrinsic motivation, and more adaptive responses to failure, whereas performance orientation, especially in its avoidance form, is associated with surface strategies, anxiety, and vulnerability to dishonesty (Anderman & Koenka, 2017; Murdock & Anderman, 2006).

Self-determination theory complements this account by distinguishing intrinsic motivation, in which activity is undertaken for its inherent satisfaction, from various forms of extrinsic regulation, and by identifying autonomy, competence, and relatedness as the basic psychological needs whose satisfaction sustains high-quality motivation (Deci & Ryan, 2000; Ryan & Deci, 2020). Empirical work in Asian higher education continues to confirm the centrality of these needs: growth needs and relatedness have been shown to significantly shape learners' motivation among Malaysian undergraduates (Hambali et al., 2025), and supportive relational environments strengthen students' capacity to withstand academic stress (Mohamed Ibrahim et al., 2025). These findings matter for the present reflection because they underline that motivation is not a fixed trait but a state cultivated or eroded by the learning environment, of which technology is now a pervasive part.

The literature on GenAI adoption suggests that current patterns of use are driven less by curiosity than by efficiency pressures. Students report using AI to save time, manage workload, and cope with deadlines (Abbas et al., 2024; Playfoot et al., 2024). Playfoot et al. (2024) found that students who feel apathetic about their degrees are more willing to use AI for assignment writing, with effort minimisation as a salient motive. This pattern points beyond the classic mastery-performance dichotomy toward what this paper terms an expedience orientation: a goal structure in which the criterion of success is neither understanding nor even demonstrated superiority, but the completion of requirements at minimal cost. Whereas performance orientation still requires the student to perform, expedience orientation delegates the performance itself. The theoretical implications of this shift are developed in the discussion.

Academic Dishonesty, Cheating, and the Integrity Tradition

Research on academic dishonesty long predates GenAI. McCabe et al. (2001), synthesising a decade of large-scale surveys, showed that cheating is widespread, that it is powerfully shaped by peer norms and institutional culture, and that honour codes and integrity climates reduce it more effectively than surveillance. Motivational research added that dishonesty is not primarily a property of bad individuals but a predictable response to goal structures: when classrooms emphasise grades and ability display, cheating rises; when they emphasise mastery and improvement, it falls (Anderman & Koenka, 2017; Murdock & Anderman, 2006). The meta-analysis by Krou et al. (2021) confirmed across more than one hundred studies that mastery orientation and intrinsic motivation are negatively associated with dishonesty, while extrinsic and avoidance orientations are positively associated with it. In short, cheating is the shadow of motivation.

The contract cheating literature is the closest precedent for AI-assisted misconduct, since it likewise involves outsourcing work to a third party (Bretag et al., 2019). Yet GenAI differs in ways that transform the moral and practical landscape: it is effectively free, instant, private, scalable, and legitimate for many uses, and it leaves no external witness. Scholars have accordingly argued that GenAI does not simply add a new cheating method but destabilises the definition of cheating itself (Cotton et al., 2024; Currie, 2023; Eaton, 2023). Eaton (2023) describes a postplagiarism era in which hybrid human-AI writing becomes normal and

historical definitions of plagiarism become inadequate, forcing institutions to rethink integrity in terms of transparency and ethical engagement rather than mere originality of text. Professional bodies have begun to respond: the European Network for Academic Integrity recommends that AI tools be acknowledged rather than banned, that outputs remain the responsibility of the human author, and that education about ethical use replace purely punitive approaches (Foltynek et al., 2023). Policy frameworks similarly emphasise developing staff and student AI literacy, clarifying acceptable use, and redesigning assessment (Chan, 2023; UNESCO, 2023).

The Contested Grey Zone Between Assistance and Dishonesty

A final strand of literature concerns the ambiguity students and staff now face. Studies of perceptions reveal wide disagreement about where legitimate use ends: using AI to explain a concept is broadly accepted, using it to polish grammar is mostly accepted, using it to generate an outline is contested, and submitting AI-generated text as one's own is broadly condemned, yet the intermediate cases form a continuum with no consensual cut point (Barrett & Pack, 2023; Chan & Hu, 2023; Yusuf et al., 2024). Fyfe (2023), reporting on assignments in which students were required to write with AI and reflect on the experience, found that students themselves struggled to articulate where their authorship began and ended. This definitional instability matters because integrity systems presuppose shared definitions of violation; when the definition dissolves, both enforcement and moral self-guidance falter (Perkins, 2023; Sullivan et al., 2023). The present reflection takes this grey zone as a central theoretical problem rather than a mere administrative nuisance.

Taken together, the literature establishes three points of departure. First, GenAI dependence is empirically real and consequential for cognition and motivation. Second, robust theory links motivation, self-regulation, and dishonesty, but this theory has not yet been systematically extended to the GenAI condition. Third, the concept of academic integrity is itself in flux. The critical reflection that follows builds on these foundations.

Cultural Variability in AI Adoption

Culture can influence people's attitudes towards technology and willingness to change established ways of working. In addition, it will also influence people using technology to decide independently or based group expectations. Perception of individual's evaluation towards AI may differ across countries (Wang, 2025). Cultural attitude also can be as a tool that can increase productivity. According to Braganza et al. (2021), they may see AI as a threat to employment and human skills. AI adoption in institutions refers to how the organization accepts, implements and regulates AI. The adoption may differ according to leadership, resources, governance and policies (Al-Marroof et al., 2025).

METHODOLOGY

Research Design

This article adopts the design of a theoretical and normative critical reflection. Critical reflection, as a scholarly genre, does not test hypotheses against new empirical data; instead, it interrogates a phenomenon by bringing established theory, existing empirical evidence, and ethical reasoning into structured dialogue in order to evaluate the phenomenon's meaning and implications. This design is appropriate for the present problem for three reasons. First, the phenomenon of student dependence on GenAI is developing faster than cumulative empirical programmes can track, so conceptual work is needed to frame the questions that empirical research should ask. Second, the core issues at stake, namely what learning means and what integrity requires, are normative questions that empirical data alone cannot settle. Third, the paper's aim is integrative: it seeks to connect literatures on self-regulated learning, achievement motivation, and academic dishonesty that have developed largely in parallel with the emerging GenAI literature. No conceptual framework in the diagrammatic sense is proposed; the strength of the paper lies in the critical argument and the integration of theory, ethics, and motivation.

Sources and Analytic Corpus

The reflection is grounded in three bodies of scholarship. The first comprises canonical and contemporary theory on motivation and self-regulation, including achievement goal theory (Ames, 1992; Dweck, 1986; Elliot & McGregor, 2001), self-determination theory (Deci & Ryan, 2000; Ryan & Deci, 2020), and self-regulated learning models (Panadero, 2017; Pintrich & De Groot, 1990; Zimmerman, 2002). The second comprises the research tradition on academic dishonesty and integrity (Anderman & Koenka, 2017; Bretag et al., 2019; Krou et al., 2021; McCabe et al., 2001; Murdock & Anderman, 2006). The third comprises the post-2022 literature on GenAI in higher education, prioritising peer-reviewed studies published between 2021 and 2026, including empirical studies of student use and its consequences (Abbas et al., 2024; Darvishi et al., 2024; Playfoot et al., 2024; Zhai et al., 2024), integrity analyses (Cotton et al., 2024; Eaton, 2023; Perkins, 2023), and policy and pedagogy frameworks (Chan, 2023; Foltynnek et al., 2023; UNESCO, 2023).

Analytic Procedure

The analysis proceeded in three steps. First, the literatures were read against one another to identify points at which the GenAI condition strains existing theory, for example where the mastery-performance dichotomy fails to describe AI-mediated behaviour. Second, from these strains, theoretical propositions were formulated: normative and explanatory claims that organise the reflection and are, in principle, empirically examinable by future research. Third, each proposition was subjected to critical evaluation in the discussion, including consideration of counterarguments and boundary conditions, before implications for practice and policy were derived. This procedure follows the accepted logic of conceptual scholarship, in which rigour is achieved through transparency of reasoning, fidelity to the cited evidence, and openness to refutation rather than through sampling and measurement.

Trustworthiness and Ethical Considerations

Because the paper analyses published literature and advances arguments rather than collecting data from human participants, no ethical clearance was required. Trustworthiness is pursued through comprehensive engagement with the relevant literatures, explicit acknowledgement of counterevidence and rival interpretations, and the clear separation of empirical claims, which are cited, from normative claims, which are argued. The authors acknowledge that any critical reflection is positioned; the position taken here is that of educators committed to the developmental purposes of higher education, and readers are invited to test the arguments against their own contexts.

Critical Reflection: Five Propositions

The integration of the three literatures yields five propositions that together describe how student dependence on GenAI reconfigures the relationships among learning, motivation, and integrity. Table 1 summarises the propositions, their theoretical anchors, and the central risk each identifies; each proposition is then elaborated and critically evaluated in the discussion.

Table 1. Theoretical propositions on AI dependence, motivation, and integrity

Proposition	Theoretical anchor	Central risk identified	Implication
P1. Learning is constituted by process, not artefact; AI dependence severs artefact from process and thereby empties task completion of learning.	Self-regulated learning theory; cognitive offloading research	Credentialed non-learning: grades certify work that produced no durable understanding.	Excessive dependence on AI may separate the final product from the learning process, allowing students to complete assignments without fully engaging in the cognitive activities necessary for meaningful learning.
P2. AI dependence shifts motivational orientation	Achievement goal theory; self-	An expedience orientation in which minimal-cost	AI dependence may shift students' motivation from learning for

from mastery toward performance and, beyond performance, toward expedience.	determination theory	completion displaces both understanding and demonstrated competence as the criterion of success.	understanding and skill development (mastery), to achieving good performance, and ultimately toward completing academic tasks as quickly and easily as possible (expedience).
P3. Habitual delegation of planning, monitoring, and evaluation to AI atrophies the self-regulatory and metacognitive capacities that higher education exists to build.	SRL models; hybrid human-AI regulation research	Metacognitive atrophy and loss of epistemic agency; skills fail to transfer once AI support is removed.	Excessive delegation of planning, monitoring, and evaluation to AI may undermine the development of students' self-regulatory and metacognitive capacities by reducing their active engagement in these cognitive processes.
P4. GenAI dissolves the shared definition of cheating, moving integrity from rule compliance to internal disposition.	Academic dishonesty research; postplagiarism scholarship	A normative grey zone in which self-serving rationalisation flourishes and enforcement loses legitimacy.	Academic integrity increasingly requires students to exercise ethical judgment and internal responsibility rather than relying solely on compliance with institutional rules.
P5. Because dishonesty is the shadow of motivation, integrity in the AI era must be cultivated through motivational and pedagogical design, not surveillance.	Motivation-cheating research; integrity culture research	Institutional over-investment in detection that punishes symptoms while classroom goal structures continue to produce the disease.	The goal should not simply be to prevent students from cheating; it should be to create learning environments in which students have meaningful reasons to learn honestly.

DISCUSSION

What Does It Mean to Learn When AI Does the Work?

The first proposition holds that learning is constituted by process rather than artefact. This claim, though simple, carries the weight of the entire reflection, so it merits careful elaboration. In the pre-AI university, artefact and process were bound together by practical necessity: a student could not normally produce a competent essay without reading, thinking, drafting, and revising, so the artefact functioned as reliable evidence of the process. Assessment systems were built on this binding. GenAI severs it. A fluent, well-structured essay can now be produced by a student who has neither read the sources nor formed a view about the question (Cotton et al., 2024; Nikolic et al., 2023). The artefact survives; the process disappears.

Self-regulated learning theory explains why this matters. In the SRL cycle, forethought, performance, and self-reflection are the engines of durable learning: goal setting activates prior knowledge, strategic effort builds and reorganises schemas, monitoring generates the feedback that calibrates self-efficacy, and reflection converts experience into transferable strategy (Panadero, 2017; Zimmerman, 2002). Every one of these mechanisms is triggered by the effortful doing of the task. Retrieval, organisation, and elaboration during writing are not incidental costs of producing text; they are the learning. When a student prompts an AI and receives finished text, none of these mechanisms fires. The student may read the output, and reading is not nothing, but passive review of a solution is among the weakest known learning activities compared with generation and retrieval. The emerging empirical record is consistent with this theoretical prediction: over-reliance on AI dialogue systems is associated with weakened critical thinking and analysis (Zhai et al., 2024), heavier GenAI use predicts poorer academic performance and greater procrastination (Abbas et al., 2024), and AI assistance produces gains that vanish when the assistance is withdrawn (Darvishi et al., 2024).

It is essential, however, to state the counterargument at full strength, because the proposition is easily caricatured as technophobia. AI can, in principle, enrich process rather than replace it. A student who uses AI to obtain alternative explanations of a difficult concept, to generate practice questions, to receive

feedback on a draft that the student then revises, or to challenge the student's argument with objections, is engaging in precisely the elaborative, feedback-rich activity that SRL theory prizes (Stojanov, 2023; Tlili et al., 2023). Molenaar (2022) and Järvelä et al. (2023) articulate this as the design question of hybrid regulation: the issue is not whether AI participates in learning but which functions it takes over and which remain with the learner. The normative line suggested by theory is that AI should support the learner's execution of the SRL cycle, not execute the cycle on the learner's behalf. The problem named by Proposition 1 is therefore not AI use as such but a specific pattern of use, namely delegation of the constitutive process, and the empirical literature indicates that this pattern is common precisely because the incentives of assessment reward artefacts, not processes (Playfoot et al., 2024). What it means to learn in the AI era must accordingly be answered in process terms: a student has learned to the extent that the student, not the tool, performed the cognitive operations the task was designed to elicit. Any conception of integrity that ignores this definition will defend the wrong thing, protecting the originality of artefacts while the learning they were meant to evidence quietly disappears.

From Mastery to Performance to Expedience: The Reordering of Motivation

The second proposition concerns motivational displacement. Achievement goal theory has long organised student motivation around the mastery-performance distinction (Ames, 1992; Elliot & McGregor, 2001). Mastery-oriented students define success as understanding; performance-oriented students define it as favourable judgement, principally grades. The dichotomy presupposes, however, that in either case the student must do the work: the mastery student to learn, the performance student to display. GenAI breaks this presupposition. When the display itself can be generated, a third orientation becomes psychologically available, one in which success is defined as discharging the requirement at minimal cost of time and effort. This paper terms it the expedience orientation.

The expedience orientation is not merely an extreme form of performance orientation, and the distinction is theoretically important. A performance-oriented student cares about the grade as a signal of ability and therefore retains some investment in the quality of the work and some anxiety about competence. An expedience-oriented student is indifferent to what the grade signals; the grade is a token to be collected so that attention can return elsewhere. Empirically, the signature of this orientation is already visible: students report efficiency and workload management as primary motives for AI use (Abbas et al., 2024; Chan & Hu, 2023), and degree apathy predicts willingness to have AI write assignments (Playfoot et al., 2024). Self-determination theory illuminates the dynamic. Deci and Ryan (2000) describe a continuum from intrinsic motivation through identified and introjected regulation to external regulation and amotivation. Expedient AI use sits at the external end and drifts toward amotivation: the activity has lost personal meaning, and only the contingency of the credential sustains compliance. Worse, the drift is self-reinforcing through the three basic needs. Competence satisfaction requires optimal challenge mastered through effort; delegating the challenge removes the experience of competence, and Bandura's (1997) mastery-experience mechanism ensures that self-efficacy stagnates. Autonomy satisfaction requires volitional engagement; a task one merely processes through a tool is experienced as alien. As need satisfaction declines, intrinsic motivation declines further, making expedient delegation still more attractive (Ryan & Deci, 2020). The result is a motivational spiral in which each act of dependence makes the next more likely.

Recent regional evidence sharpens the concern. Rifin et al. found strong positive links between self-regulated strategy use and both intrinsic motivation and self-efficacy among Malaysian students, implying that anything that suppresses strategy use, as wholesale AI delegation does, can be expected to depress the motivational beliefs that sustain learning (Rifin et al., 2025). Likewise, evidence that growth and relatedness needs drive learner motivation (Hambali et al., 2025) suggests a further mechanism: a student who works primarily with a machine interlocutor, rather than with teachers and peers, forgoes the relational nourishment that keeps motivation alive. The counterargument must again be heard: for some students AI lowers barriers, reduces anxiety, and provides the early success experiences that build efficacy, particularly for second language writers and students with weak preparation (Chan & Hu, 2023; Sullivan et al., 2023). The proposition therefore does not claim that AI use causes expedience in all users; it claims that AI availability makes expedience a stable, low-cost attractor in the motivational landscape, and that assessment

regimes rewarding artefacts over processes will steer increasing numbers of students into it. The practical corollary is that motivational design, long treated as a soft concern, becomes the hard core of academic integrity strategy.

Outsourcing the Self: Metacognitive Offloading and the Atrophy of Self-Regulation

The third proposition extends the analysis from motivation to capability. SRL is not only a description of how people learn; it is itself a developmental achievement, a set of capacities in planning, monitoring, evaluating, and strategy selection that students acquire through repeated, effortful practice (Panadero, 2017; Zimmerman, 2002). Higher education has traditionally been an apprenticeship in these capacities: the unstructured essay, the long project, and the open problem force students to decide what to do, notice when it is not working, and adjust. GenAI now offers to perform exactly these regulatory functions. It will plan the essay, structure the argument, evaluate the draft, diagnose the error, and schedule the study plan. The tool does not merely offload cognition; it offloads metacognition.

The cognitive offloading literature shows that people who expect information to be externally available invest less in encoding it (Risko & Gilbert, 2016; Sparrow et al., 2011). Applied to metacognition, the prediction is that students who expect AI to monitor and evaluate for them will invest less in developing internal standards of quality, and there is early evidence pointing this way: negative associations between AI reliance and critical thinking (Gerlich, 2025; Zhai et al., 2024), reduced self-reported analytic effort among confident AI users (Lee et al., 2025), and the failure of AI-supported gains to persist without the tool (Darvishi et al., 2024). The deepest cost here is what may be called epistemic agency: the capacity and disposition to form, test, and take responsibility for one's own judgements. A graduate without internalised standards cannot evaluate AI output either, which produces the paradox that heavy AI dependence in education undermines the very competence needed to use AI well in professional life (Bearman & Ajjawi, 2023; Markauskaite et al., 2022).

The boundary condition on this proposition concerns expertise. For experts, offloading routine operations is rational and even virtuous; it frees capacity for higher-order work. The developmental problem is specific to novices, who are asked to offload operations they have not yet acquired. This suggests a principle of pedagogical sequencing: AI delegation should expand with demonstrated competence, so that students first build the capacity and only then decide what to delegate (Järvelä et al., 2023; Molenaar, 2022). It also suggests that blanket institutional answers, whether prohibition or laissez-faire, are both wrong, because the developmental meaning of the same act of delegation differs across levels of expertise. Resilience research offers a complementary insight: capacities for coping with academic difficulty are built through supported exposure to stressors, not through their removal (Mohamed Ibrahim et al., 2025). An educational environment in which AI removes all friction is, in this sense, an environment that manufactures fragility.

Rethinking Cheating: Integrity in the Grey Zone

The fourth proposition addresses integrity directly. Traditional academic integrity regimes rest on a tripod: a shared definition of violation, a detection capability, and a sanction system. GenAI weakens all three legs simultaneously. Definition falters because AI assistance forms a continuum from spell-checking to full generation with no natural cut point, and stakeholders demonstrably disagree about where legitimacy ends (Barrett & Pack, 2023; Chan & Hu, 2023; Yusuf et al., 2024). Detection falters because AI text bears no stable fingerprint; detection tools produce both false negatives and, more corrosively, false accusations (Cotton et al., 2024; Perkins, 2023). Sanctioning falters because enforcement without reliable definition and detection is experienced as arbitrary, which drains the legitimacy on which compliance depends (Currie, 2023; Sullivan et al., 2023).

Eaton's (2023) postplagiarism thesis draws the radical conclusion: in an era when hybrid human-AI writing becomes the norm, integrity can no longer be defined as textual originality. What, then, should it be defined as? The literature converges on two elements. The first is transparency: honest disclosure of how AI was used, so that assessors can judge what the artefact evidences (Foltynek et al., 2023; UNESCO, 2023). The second, argued here to be more fundamental, is fidelity to the educational purpose of the task: a use of AI is

legitimate insofar as it leaves intact the cognitive process the task was designed to elicit, and illegitimate insofar as it substitutes for that process. This purposive definition explains the intuitions that the continuum view leaves mysterious. Using AI to explain a concept is legitimate because comprehension is aided, not replaced; having AI write the analysis is illegitimate in an analysis task because the assessed process has been outsourced, exactly as in contract cheating (Bretag et al., 2019). The definition also implies task relativity: the same use of AI may be legitimate in one task and dishonest in another, depending on what the task is for. This is uncomfortable for rule-writers but faithful to the moral reality.

The deeper implication is a shift in the locus of integrity. When violations were externally definable and detectable, integrity could be practically treated as compliance. In the grey zone, where the decisive facts about process are known only to the student, integrity necessarily becomes an internal disposition: the settled preference for honest learning even when dishonesty is safe. Moral psychology warns how demanding this is; rationalisation flourishes precisely where definitions are contested, and students can tell themselves that everyone uses it, that the tool is legal, that only the output matters (McCabe et al., 2001; Murdock & Anderman, 2006). This is why the present paper insists that the integrity question and the motivation question are one question. A student with mastery goals has an internal reason not to delegate the process, because the process is the point; a student with expedience goals has none. Institutions that respond to the grey zone with detection arms races are therefore treating the symptom. The disease is motivational, and it is to institutional responsibility for that disease that the final proposition turns.

Character, Virtue, and the Long-Term Formation of the Learner

Before institutional implications, the reflection must confront its most normative register: what repeated expedient delegation does to the person. Virtue ethics holds that character is formed by habituation; we become what we repeatedly do. The intellectual virtues relevant to academic life, among them perseverance, honesty, humility, curiosity, and intellectual courage, are acquired through repeated acts of persevering, truth-telling, and inquiring, especially when these are costly. Each assignment is, in this light, not merely an assessment event but a rehearsal of character. A student who habitually chooses effortful engagement rehearses perseverance and honesty; a student who habitually delegates rehearses avoidance and misrepresentation. Over a degree programme comprising hundreds of such rehearsals, the difference compounds into two different kinds of person (Crawford et al., 2023).

The long-term stakes extend beyond the individual. Higher education's social warrant is the claim that its graduates possess formed judgement; professions, employers, and democratic publics rely on that warrant. If AI dependence produces graduates whose credentials certify processes that never occurred, the warrant is falsified at scale, with consequences for professional competence and public trust that will surface years after the assessments in question (Currie, 2023; Dwivedi et al., 2023). There is also an equity dimension aligned with SDG 4: students with strong prior preparation and educated support networks are more likely to use AI as an amplifier of learning, while students under greater time and economic pressure face stronger temptations toward expedient substitution (Abbas et al., 2024). Unmanaged, GenAI may thus widen precisely the quality gaps that SDG 4 commits the global community to closing, even while headline participation metrics improve (UNESCO, 2023). Quality education, in the meaning of SDG 4, is inseparable from the formation of capable, honest, self-regulating learners; a system that certifies without forming has abandoned the goal in substance while honouring it in statistics.

Implications for Pedagogy, Assessment, and Institutional Policy

The fifth proposition, that integrity must be cultivated through motivational and pedagogical design rather than surveillance, yields concrete directions. For pedagogy, the priority is to make process visible and valuable. Assessment should sample the process directly through drafts, oral defences, in-class writing, project logs, and vivas, so that the artefact ceases to be the sole evidence of learning (Nikolic et al., 2023; Rudolph et al., 2023). Structured AI-transparency statements, in which students disclose and reflect on their AI use, convert the grey zone into an object of learning and rehearse the honesty the era requires (Fyfe, 2023; Foltynnek et al., 2023). Teaching students explicitly about SRL, metacognition, and the process-

dependence of learning gives them the reason to choose effort that prohibition cannot supply (Bearman & Ajjaw, 2023; Zimmerman, 2002).

For motivational design, achievement goal research prescribes mastery-structured environments: emphasis on improvement and understanding, tolerance of error, meaningful tasks connected to students' purposes, and de-emphasis of comparative grading, all of which are associated with reduced dishonesty (Ames, 1992; Anderman & Koenka, 2017; Krou et al., 2021). Self-determination theory adds the need-supportive triad of autonomy support, optimal challenge, and relational warmth (Ryan & Deci, 2020), and regional evidence confirms that relatedness and growth remain potent levers for Asian learners (Hambali et al., 2025). Authentic assessments that matter to real audiences give students reasons to want the competence rather than merely the credential.

For institutional policy, the argument counsels a decisive rebalancing from detection toward development. Policies should define legitimate use purposively and task-relatively, require transparency rather than promising detection, invest in staff and student AI literacy, and treat first instances of grey-zone misuse educatively rather than punitively (Chan, 2023; Eaton, 2023; Perkins, 2023; UNESCO, 2023). Integrity climate research shows that shared norms, honour commitments, and visible institutional seriousness about learning reduce dishonesty more reliably than surveillance (McCabe et al., 2001). None of this implies abandoning sanctions for clear fraud; it implies recognising that in the grey zone the decisive battle is for students' motivational orientation, and that battle is won or lost in curriculum and classroom design, not in detection software.

Counterarguments and Boundary Conditions

Three counterarguments deserve explicit answers. The first holds that history refutes alarm: calculators, word processors, and search engines each provoked similar fears, and education adapted. The reply is that the analogy understates the novelty. Previous tools automated operations peripheral to the assessed competence; GenAI automates the assessed competence itself, including the metacognitive functions that previous tools left untouched (Selwyn, 2022; Zhai et al., 2024). Adaptation is indeed possible, but it will require the deliberate redesign argued for here, not mere passage of time. The second counterargument holds that AI literacy is now the real learning outcome, so delegation is not avoidance but the exercise of a new competence. The reply is partial concession: prompting, evaluating, and integrating AI output are genuine and important capabilities (Markauskaite et al., 2022). But their exercise presupposes the internal standards that only non-delegated learning builds; one cannot critically evaluate an analysis one could never have produced. AI literacy is therefore a supplement to formed competence, not a substitute for its formation. The third counterargument holds that the expedience diagnosis blames students for a rational response to credentialist systems that universities themselves built. This objection is accepted rather than rebutted: it is the core of Proposition 5. Students' expedience is the mirror of institutional goal structures, and the responsibility for reform lies chiefly with institutions (Murdock & Anderman, 2006; Playfoot et al., 2024). The critical reflection is directed at systems before it is directed at students.

CONCLUSION

This article has offered a critical reflection on academic motivation and integrity in the era of generative AI, organised around five propositions. Learning, it argued, is constituted by cognitive process rather than by artefacts, so that habitual delegation of tasks to AI empties completion of learning even where no rule is broken. Dependence on AI reshapes motivational orientation, adding to the classic mastery and performance goals an expedience orientation in which minimal-cost completion becomes the criterion of success, sustained by a self-reinforcing spiral of declining need satisfaction and self-efficacy. Delegation of planning, monitoring, and evaluation to AI risks atrophying the self-regulatory and metacognitive capacities that constitute epistemic agency, particularly for novices. GenAI dissolves shared definitions of cheating, relocating integrity from external compliance to internal disposition and rendering purposive, task-relative definitions of legitimate use necessary. Finally, because dishonesty is the shadow of motivation, the

appropriate institutional response is motivational and pedagogical redesign, with transparency and mastery-oriented, process-visible assessment at its centre, rather than an arms race of detection.

The contribution of the paper is integrative and normative. It connects the post-2022 GenAI literature to the mature research traditions on self-regulated learning, achievement motivation, and academic dishonesty, and it derives from that connection a vocabulary, most notably the expedience orientation and the purposive definition of legitimate AI use, that can guide both research and policy. Practically, the paper directs institutions toward assessment that samples process, policies that require transparency, and classroom climates that make mastery worth wanting, positioning these not as pedagogical niceties but as the substantive core of integrity strategy and as conditions for delivering the quality education promised by SDG 4.

The limitations of the work follow from its genre. As a conceptual reflection it establishes plausibility and coherence, not causal fact; its propositions are offered as examinable claims. The empirical base on which it draws is young, dominated by self-report and cross-sectional designs, and evolving with the technology itself. The reflection is also written from a higher education standpoint and, in part, from an Asian institutional context; its arguments should be tested against other levels and cultures of education. Future research should trace motivational orientations longitudinally among students with differing patterns of AI use, operationalise and validate the expedience construct, test whether process-visible assessment and mastery climates reduce illegitimate delegation, and examine how transparency requirements shape both honesty and learning. The question these studies must ultimately answer is the one this reflection has tried to sharpen: how education can form people who choose effortful learning when effortless output is always at hand. On that choice, and on institutions' willingness to make it a reasonable one, the integrity of higher education in the AI era now depends.

ACKNOWLEDGEMENTS

The authors would like to thank UNITAR International University for supporting this research.

REFERENCES

1. Al-Marooof, R. S., et al. (2025). A meta-analysis of TOE factors driving organizational adoption of artificial intelligence across industries. *Discover Artificial Intelligence*. <https://doi.org/10.1007/s44163-025-00747-2>
2. Abbas, M., Jam, F. A., & Khan, T. I. (2024). Is it harmful or helpful? Examining the causes and consequences of generative AI usage among university students. *International Journal of Educational Technology in Higher Education*, 21, 10. <https://doi.org/10.1186/s41239-024-00444-7>
3. Adiguzel, T., Kaya, M. H., & Cansu, F. K. (2023). Revolutionizing education with AI: Exploring the transformative potential of ChatGPT. *Contemporary Educational Technology*, 15(3), ep429. <https://doi.org/10.30935/cedtech/13152>
4. Ames, C. (1992). Classrooms: Goals, structures, and student motivation. *Journal of Educational Psychology*, 84(3), 261-271. <https://doi.org/10.1037/0022-0663.84.3.261>
5. Anderman, E. M., & Koenka, A. C. (2017). The relation between academic motivation and cheating. *Theory Into Practice*, 56(2), 95-102. <https://doi.org/10.1080/00405841.2017.1308172>
6. Bandura, A. (1997). *Self-efficacy: The exercise of control*. W. H. Freeman.
7. Barrett, A., & Pack, A. (2023). Not quite eye to A.I.: Student and teacher perspectives on the use of generative artificial intelligence in the writing process. *International Journal of Educational Technology in Higher Education*, 20, 59. <https://doi.org/10.1186/s41239-023-00427-0>
8. Bearman, M., & Ajjawi, R. (2023). Learning to work with the black box: Pedagogy for a world with artificial intelligence. *British Journal of Educational Technology*, 54(5), 1160-1173. <https://doi.org/10.1111/bjet.13337>
9. Bozkurt, A., Xiao, J., Lambert, S., Pazurek, A., Crompton, H., Koseoglu, S., & Jandric, P. (2023). Speculative futures on ChatGPT and generative artificial intelligence (AI): A collective reflection

- from the educational landscape. *Asian Journal of Distance Education*, 18(1), 53-130. <https://doi.org/10.5281/zenodo.7636568>
10. Braganza, A., Chen, W., Canhoto, A., & Sap, S. (2021). Productive employment and decent work: The impact of AI adoption on psychological contracts, job engagement and employee trust. *Journal of Business Research*, 131, 485–494. <https://doi.org/10.1016/j.jbusres.2020.08.018>
11. Bretag, T., Harper, R., Burton, M., Ellis, C., Newton, P., Rozenberg, P., Saddiqui, S., & van Haeringen, K. (2019). Contract cheating: A survey of Australian university students. *Studies in Higher Education*, 44(11), 1837-1856. <https://doi.org/10.1080/03075079.2018.1462788>
12. Chan, C. K. Y. (2023). A comprehensive AI policy education framework for university teaching and learning. *International Journal of Educational Technology in Higher Education*, 20, 38. <https://doi.org/10.1186/s41239-023-00408-3>
13. Chan, C. K. Y., & Hu, W. (2023). Students' voices on generative AI: Perceptions, benefits, and challenges in higher education. *International Journal of Educational Technology in Higher Education*, 20, 43. <https://doi.org/10.1186/s41239-023-00411-8>
14. Cotton, D. R. E., Cotton, P. A., & Shipway, J. R. (2024). Chatting and cheating: Ensuring academic integrity in the era of ChatGPT. *Innovations in Education and Teaching International*, 61(2), 228-239. <https://doi.org/10.1080/14703297.2023.2190148>
15. Crawford, J., Cowling, M., & Allen, K.-A. (2023). Leadership is needed for ethical ChatGPT: Character, assessment, and learning using artificial intelligence (AI). *Journal of University Teaching and Learning Practice*, 20(3). <https://doi.org/10.53761/1.20.3.02>
16. Currie, G. M. (2023). Academic integrity and artificial intelligence: Is ChatGPT hype, hero or heresy? *Seminars in Nuclear Medicine*, 53(5), 719-730. <https://doi.org/10.1053/j.semnuclmed.2023.04.008>
17. Darvishi, A., Khosravi, H., Sadiq, S., Gasevic, D., & Siemens, G. (2024). Impact of AI assistance on student agency. *Computers & Education*, 210, 104967. <https://doi.org/10.1016/j.compedu.2023.104967>
18. Deci, E. L., & Ryan, R. M. (2000). The “what” and “why” of goal pursuits: Human needs and the self-determination of behavior. *Psychological Inquiry*, 11(4), 227-268. https://doi.org/10.1207/S15327965PLI1104_01
19. Dweck, C. S. (1986). Motivational processes affecting learning. *American Psychologist*, 41(10), 1040-1048. <https://doi.org/10.1037/0003-066X.41.10.1040>
20. Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., Baabdullah, A. M., Koohang, A., Raghavan, V., Ahuja, M., & Wright, R. (2023). “So what if ChatGPT wrote it?” Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *International Journal of Information Management*, 71, 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
21. Eaton, S. E. (2023). Postplagiarism: Transdisciplinary ethics and integrity in the age of artificial intelligence and neurotechnology. *International Journal for Educational Integrity*, 19, 23. <https://doi.org/10.1007/s40979-023-00144-1>
22. Elliot, A. J., & McGregor, H. A. (2001). A 2 x 2 achievement goal framework. *Journal of Personality and Social Psychology*, 80(3), 501-519. <https://doi.org/10.1037/0022-3514.80.3.501>
23. Farrelly, T., & Baker, N. (2023). Generative artificial intelligence: Implications and considerations for higher education practice. *Education Sciences*, 13(11), 1109. <https://doi.org/10.3390/educsci13111109>
24. Foltynnek, T., Bjelobaba, S., Glendinning, I., Khan, Z. R., Santos, R., Pavletic, P., & Kravjar, J. (2023). ENAI recommendations on the ethical use of artificial intelligence in education. *International Journal for Educational Integrity*, 19, 12. <https://doi.org/10.1007/s40979-023-00133-4>
25. Fyfe, P. (2023). How to cheat on your final paper: Assigning AI for student writing. *AI & Society*, 38, 1395-1405. <https://doi.org/10.1007/s00146-022-01397-z>
26. Gerlich, M. (2025). AI tools in society: Impacts on cognitive offloading and the future of critical thinking. *Societies*, 15(1), 6. <https://doi.org/10.3390/soc15010006>
27. Hambali, N., Muhamad, S. H. A., Noh, N., Alias, A. F. N., & Rahmat, N. H. (2025). Exploring the influence of existence, relatedness and growth on learner's motivation. *Asian Journal of Social Science Research*, 7(1), 104-124. <https://doi.org/10.5281/zenodo.15804199>

28. Holmes, W., & Tuomi, I. (2022). State of the art and practice in AI in education. *European Journal of Education*, 57(3), 542-570. <https://doi.org/10.1111/ejed.12533>
29. Jarvela, S., Nguyen, A., & Hadwin, A. (2023). Human and artificial intelligence collaboration for socially shared regulation in learning. *British Journal of Educational Technology*, 54(5), 1057-1076. <https://doi.org/10.1111/bjet.13325>
30. Kasneci, E., Sessler, K., Kuchemann, S., Bannert, M., Dementieva, D., Fischer, F., Gasser, U., Groh, G., Gunnemann, S., Hullermeier, E., Krusche, S., Kutyniok, G., Michaeli, T., Nerdel, C., Pfeffer, J., Poquet, O., Sailer, M., Schmidt, A., Seidel, T., & Kasneci, G. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103, 102274. <https://doi.org/10.1016/j.lindif.2023.102274>
31. Krou, M. R., Fong, C. J., & Hoff, M. A. (2021). Achievement motivation and academic dishonesty: A meta-analytic investigation. *Educational Psychology Review*, 33, 427-458. <https://doi.org/10.1007/s10648-020-09557-7>
32. Lee, H.-P., Sarkar, A., Tankelevitch, L., Drosos, I., Rintel, S., Banks, R., & Wilson, N. (2025). The impact of generative AI on critical thinking: Self-reported reductions in cognitive effort and confidence effects from a survey of knowledge workers. *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*. ACM. <https://doi.org/10.1145/3706598.3713778>
33. Lo, C. K. (2023). What is the impact of ChatGPT on education? A rapid review of the literature. *Education Sciences*, 13(4), 410. <https://doi.org/10.3390/educsci13040410>
34. Markauskaite, L., Marrone, R., Poquet, O., Knight, S., Martinez-Maldonado, R., Howard, S., Tondeur, J., De Laat, M., Buckingham Shum, S., Gasevic, D., & Siemens, G. (2022). Rethinking the entwinement between artificial intelligence and human learning: What capabilities do learners need for a world with AI? *Computers and Education: Artificial Intelligence*, 3, 100056. <https://doi.org/10.1016/j.caeai.2022.100056>
35. McCabe, D. L., Trevino, L. K., & Butterfield, K. D. (2001). Cheating in academic institutions: A decade of research. *Ethics & Behavior*, 11(3), 219-232. https://doi.org/10.1207/S15327019EB1103_2
36. Memarian, B., & Doleck, T. (2023). Fairness, accountability, transparency, and ethics (FATE) in artificial intelligence (AI) and higher education: A systematic review. *Computers and Education: Artificial Intelligence*, 5, 100152. <https://doi.org/10.1016/j.caeai.2023.100152>
37. Michel-Villareal, R., Vilalta-Perdomo, E., Salinas-Navarro, D. E., Thierry-Aguilera, R., & Gerardou, F. S. (2023). Challenges and opportunities of generative AI for higher education as explained by ChatGPT. *Education Sciences*, 13(9), 856. <https://doi.org/10.3390/educsci13090856>
38. Mohamed Ibrahim, S. B., Mohd Basir, S. N., Johanis, M. A., Habib Sultan, N. H., & Joharis, N. F. (2025). Exploring the impact of social support and academic stress on engineering students' resiliency. *Asian Journal of Social Science Research*, 7(1), 66-80. <https://doi.org/10.5281/zenodo.15803407>
39. Molenaar, I. (2022). Towards hybrid human-AI learning technologies. *European Journal of Education*, 57(4), 632-645. <https://doi.org/10.1111/ejed.12527>
40. Murdock, T. B., & Anderman, E. M. (2006). Motivational perspectives on student cheating: Toward an integrated model of academic dishonesty. *Educational Psychologist*, 41(3), 129-145. https://doi.org/10.1207/s15326985ep4103_1
41. Ngo, T. T. A. (2023). The perception by university students of the use of ChatGPT in education. *International Journal of Emerging Technologies in Learning*, 18(17), 4-19. <https://doi.org/10.3991/ijet.v18i17.39019>
42. Nikolic, S., Daniel, S., Haque, R., Belkina, M., Hassan, G. M., Grundy, S., Lyden, S., Neal, P., & Sandison, C. (2023). ChatGPT versus engineering education assessment: A multidisciplinary and multi-institutional benchmarking and analysis of this generative artificial intelligence tool to investigate assessment integrity. *European Journal of Engineering Education*, 48(4), 559-614. <https://doi.org/10.1080/03043797.2023.2213169>
43. Panadero, E. (2017). A review of self-regulated learning: Six models and four directions for research. *Frontiers in Psychology*, 8, 422. <https://doi.org/10.3389/fpsyg.2017.00422>

44. Perkins, M. (2023). Academic integrity considerations of AI large language models in the post-pandemic era: ChatGPT and beyond. *Journal of University Teaching and Learning Practice*, 20(2). <https://doi.org/10.53761/1.20.02.07>
45. Pintrich, P. R., & De Groot, E. V. (1990). Motivational and self-regulated learning components of classroom academic performance. *Journal of Educational Psychology*, 82(1), 33-40. <https://doi.org/10.1037/0022-0663.82.1.33>
46. Playfoot, D., Quigley, M., & Thomas, A. G. (2024). Hey ChatGPT, give me a title for a paper about degree apathy and student use of AI for assignment writing. *The Internet and Higher Education*, 62, 100950. <https://doi.org/10.1016/j.iheduc.2024.100950>
47. Rasul, T., Nair, S., Kalendra, D., Robin, M., de Oliveira Santini, F., Ladeira, W. J., Sun, M., Day, I., Rather, R. A., & Heathcote, L. (2023). The role of ChatGPT in higher education: Benefits, challenges, and future research directions. *Journal of Applied Learning and Teaching*, 6(1), 41-56. <https://doi.org/10.37074/jalt.2023.6.1.29>
48. Rifin, R., Sidik, N. B., Abdul Rahman, N. H., Mohd Hussain, M., & Mokhtar, S. M. A. (2025). The influence of self-regulated strategies on components in motivational beliefs. *Asian Journal of Social Science Research*, 7(1), 146-160. <https://doi.org/10.5281/zenodo.15804446>
49. Risko, E. F., & Gilbert, S. J. (2016). Cognitive offloading. *Trends in Cognitive Sciences*, 20(9), 676-688. <https://doi.org/10.1016/j.tics.2016.07.002>
50. Rudolph, J., Tan, S., & Tan, S. (2023). ChatGPT: Bullshit spewer or the end of traditional assessments in higher education? *Journal of Applied Learning and Teaching*, 6(1), 342-363. <https://doi.org/10.37074/jalt.2023.6.1.9>
51. Ryan, R. M., & Deci, E. L. (2020). Intrinsic and extrinsic motivation from a self-determination theory perspective: Definitions, theory, practices, and future directions. *Contemporary Educational Psychology*, 61, 101860. <https://doi.org/10.1016/j.cedpsych.2020.101860>
52. Selwyn, N. (2022). The future of AI and education: Some cautionary notes. *European Journal of Education*, 57(4), 620-631. <https://doi.org/10.1111/ejed.12532>
53. Sparrow, B., Liu, J., & Wegner, D. M. (2011). Google effects on memory: Cognitive consequences of having information at our fingertips. *Science*, 333(6043), 776-778. <https://doi.org/10.1126/science.1207745>
54. Stojanov, A. (2023). Learning with ChatGPT 3.5 as a more knowledgeable other: An autoethnographic study. *International Journal of Educational Technology in Higher Education*, 20, 35. <https://doi.org/10.1186/s41239-023-00404-7>
55. Sullivan, M., Kelly, A., & McLaughlan, P. (2023). ChatGPT in higher education: Considerations for academic integrity and student learning. *Journal of Applied Learning and Teaching*, 6(1), 31-40. <https://doi.org/10.37074/jalt.2023.6.1.17>
56. Tlili, A., Shehata, B., Adarkwah, M. A., Bozkurt, A., Hickey, D. T., Huang, R., & Agyemang, B. (2023). What if the devil is my guardian angel: ChatGPT as a case study of using chatbots in education. *Smart Learning Environments*, 10, 15. <https://doi.org/10.1186/s40561-023-00237-x>
57. UNESCO. (2023). Guidance for generative AI in education and research. UNESCO Publishing. <https://doi.org/10.54675/EWZM9535>
58. Yusuf, A., Pervin, N., & Roman-Gonzalez, M. (2024). Generative AI and the future of higher education: A threat to academic integrity or reformation? Evidence from multicultural perspectives. *International Journal of Educational Technology in Higher Education*, 21, 21. <https://doi.org/10.1186/s41239-024-00453-6>
59. Wang, S. (2025). Public perceptions of artificial intelligence in 20 countries: Assessing individual- and country-level factors. *Cross-Cultural Research*, 59(5), 651-676. <https://doi.org/10.1177/10693971251336803>
60. Zhai, C., Wibowo, S., & Li, L. D. (2024). The effects of over-reliance on AI dialogue systems on students' cognitive abilities: A systematic review. *Smart Learning Environments*, 11, 28. <https://doi.org/10.1186/s40561-024-00316-7>
61. Zimmerman, B. J. (2002). Becoming a self-regulated learner: An overview. *Theory Into Practice*, 41(2), 64-70. https://doi.org/10.1207/s15430421tip4102_2