



REVIEW OF ARTIFICIAL INTELLIGENCE IN EDUCATION

DOI: <https://doi.org/10.37497/rev.artif.intell.educ.v7ii.1381>

ARTICLE

MODELLING HUMAN GENERATIVE AI INTERACTION IN HIGHER EDUCATION: EVALUATION AMONG GHANAIAAN STUDENTS

Modelagem da Interação Humano–Inteligência
Artificial Generativa no Ensino Superior:
Avaliação entre Estudantes Ganenses

Adamu Wahab

Department of Construction Technology and Management Education, University of Skills Training and Entrepreneurial
Development, Kumasi, Ghana.
E-mail: wahab.adamu3500@gmail.com

Thomas Mensah Asabere

Department of Marketing and Entrepreneurship, University of Ghana Business School, Greater Accra, Ghana.
Email: tmasabere@gmail.com

Issah Chakurah

Department of Vocational and Technical Education, University of Cape Coast, Cape Coast, Ghana.
Email: chakurah.issah@ucc.edu.gh



Received: 15 May 2026

Revised: 30 June 2026

Accepted: 12 August 2026

e-ISSN: 2965-4688

Corresponding Author: Adamu Wahab –
Email: wahab.adamu3500@gmail.com

How to cite this article: Wahab, A.,
Asabere, T. M., & Chakurah, I. (2026).
Modelling Human Generative AI
Interaction in Higher Education:
Evaluation Among Ghanaian
Students. *Review of Artificial
Intelligence in Education*, 7(i), e01381.
<https://doi.org/10.37497/rev.artif.intell.educ.v7ii.1381>

ABSTRACT | Background: Generative artificial intelligence (GenAI) is rapidly transforming higher education, yet the cognitive processes through which it affects learning remain under-theorised. **Objective:** This study extends adoption-focused perspectives by developing and empirically testing a Human-AI Interaction Model (HAIM), which positions the quality of cognitive interaction, rather than mere adoption, as the key determinant of educational outcomes. The model draws on Extended Cognition Theory and Self-Regulated Learning Theory and specifies ten hypothesised pathways among seven constructs: AI literacy, trust in AI, cognitive offloading, co-creation behaviour, verification behaviour, learning outcomes, and critical thinking. **Methods:** Cross-sectional survey data were collected from 623 university students across five Ghanaian institutions and analysed using Partial Least Squares Structural Equation Modelling (PLS-SEM) with 5,000 bootstrap resamples. The measurement model exhibited good reliability, and convergent and discriminant validity were confirmed. **Results:** AI literacy positively predicted co-creation behaviour ($\beta = .479$) and negatively predicted cognitive offloading ($\beta = -.199$). Uncalibrated trust was positively associated with cognitive offloading ($\beta = .319$), which in turn reduced learning outcomes ($\beta = -.216$). Co-creation behaviour was the strongest predictor of learning outcomes ($\beta = .531$), while verification behaviour was the strongest predictor of critical thinking ($\beta = .540$). Mediation analysis confirmed that cognitive offloading and verification behaviour are the primary mechanisms linking trust and literacy to outcomes. **Conclusion:** The findings support a cognitively informed account of human-AI interaction and underscore the need for AI literacy interventions that cultivate critical engagement rather than passive dependence, particularly within Global South higher education contexts.

Keywords | Generative AI; Human-AI interaction; Cognitive Offloading; AI literacy; Critical Thinking; PLS-SEM; Ghana; Higher Education





RESUMO | Contexto: A inteligência artificial generativa (IAGen) está transformando rapidamente o ensino superior; contudo, os processos cognitivos por meio dos quais ela afeta a aprendizagem ainda permanecem insuficientemente teorizados. **Objetivo:** Este estudo amplia as perspectivas centradas na adoção ao desenvolver e testar empiricamente um Modelo de Interação Humano-IA (HAIM), que posiciona a qualidade da interação cognitiva, e não a mera adoção da tecnologia, como o principal determinante dos resultados educacionais. O modelo fundamenta-se na Teoria da Cognição Estendida e na Teoria da Aprendizagem Autorregulada e especifica dez relações hipotéticas entre sete construtos: letramento em IA, confiança na IA, descarregamento cognitivo, comportamento de cocriação, comportamento de verificação, resultados de aprendizagem e pensamento crítico. **Métodos:** Foram coletados dados de uma pesquisa transversal com 623 estudantes universitários de cinco instituições de Gana, analisados por meio da Modelagem de Equações Estruturais por Mínimos Quadrados Parciais (PLS-SEM), com 5.000 reamostragens bootstrap. O modelo de mensuração apresentou boa confiabilidade, tendo sido confirmadas as validades convergente e discriminante. **Resultados:** O letramento em IA predisse positivamente o comportamento de cocriação ($\beta = 0,479$) e negativamente o descarregamento cognitivo ($\beta = -0,199$). A confiança não calibrada esteve positivamente associada ao descarregamento cognitivo ($\beta = 0,319$), que, por sua vez, reduziu os resultados de aprendizagem ($\beta = -0,216$). O comportamento de cocriação foi o mais forte preditor dos resultados de aprendizagem ($\beta = 0,531$), enquanto o comportamento de verificação foi o mais forte preditor do pensamento crítico ($\beta = 0,540$). A análise de mediação confirmou que o descarregamento cognitivo e o comportamento de verificação constituem os principais mecanismos que conectam a confiança e o letramento aos resultados. **Conclusão:** Os achados sustentam uma abordagem cognitivamente fundamentada da interação humano-IA e ressaltam a necessidade de intervenções voltadas ao letramento em IA que promovam o engajamento crítico, em vez da dependência passiva, particularmente nos contextos de ensino superior do Sul Global.

Palavras-chave | Inteligência Artificial Generativa; interação humano-IA; descarregamento cognitivo; letramento em IA; pensamento crítico; PLS-SEM; Gana; ensino superior

INTRODUCTION

The implementation of generative artificial intelligence (GenAI) tools within university learning settings has been a slow-moving process that has been far behind the curve of theoretical and empirical explanations of its educational implications. Student academic life around the world, in a wide range of institutional contexts, disciplinary traditions, and socioeconomic settings, has become routine and predictable through platforms built on large language models (LLMs), such as ChatGPT, Gemini, Microsoft Copilot and their variants (Alshamy et al., 2025; Farrelly and Baker, 2023; Francis et al., 2025; Granć, 2025). The very magnitude of such diffusion is consequential: questions of how students cognitively interact with AI-based tools have ceased to be speculative but urgently empirical, with real-time implications on the quality of learning, academic integrity, and developmental trajectories of millions of students in the world (Abbas et al., 2024; Ahangama, 2026; Essel et al., 2023).

Nevertheless, despite the urgency of these questions, the prevailing theoretical frameworks in the research of educational technology research are stuck in a previous technological era of adoption models. The Technology Acceptance Model (TAM) and the Unified Theory of Acceptance and Use of Technology (UTAUT) – and its many derivatives – view technology integration as a fundamentally decision-making problem: will students decide to use the tool? (Goto & Ramnarain, 2026; Hazaimah & Al-Ansi, 2024; Sergeeva et al., 2025; Weis et al., 2026). This framing proved fruitful in cognising the uptake of comparatively straightforward, task-specific digital devices. But GenAI systems are



qualitatively different to the previous educational technologies: they are generative, conversational, cognitively pervasive, and capable of substituting the cognitive processes of reasoning, drafting, and synthesising, and problem-solving, which are the very substance of higher learning. The application of adoption-specific frameworks to GenAI thus commits a category error: it quantifies the use of AI by students, when the most important question in education is how students engage cognitively with AI and the cognitive effects of this engagement (Haroud and Saqri, 2025; Kim et al., 2025; Strzelecki and ElArabawy, 2024).

The difference between adoption and the quality of interaction is of significant theoretical and practical interest. Studies that simply record rates of GenAI adoption cannot be able to differentiate between a student who uses AI as a dialogic intellectual partner, iterating prompts, interrogating outputs, synthesising AI-generated content with independent reasoning, and a student who delegates the entirety of their academic cognitive work to AI without engagement or examination. Two students can have the same profiles of adoption but very different learning outcomes, development of critical thinking, and long-term trajectories of academic competence (Ayanwale et al., 2025; Crompton et al., 2026; Wu and Zhang, 2025; Zhao et al., 2025). Educational literature is desperately in need of models that can model and explain these qualitative differences in interaction instead of collapsing them into binary variables in adoption.

This theoretical gap in the Ghanaian context of higher education, the empirical setting of the present study, has especially acute practical implications. The public universities of Ghana have seen a rapid uptake of GenAI by the student body, but no evidence of an institutional structure of AI governance, organised AI literacy education, or a grounded set of guidelines on pedagogically productive AI engagement (Asghar et al., 2025; Mante et al., 2025; Wiredu et al., 2024). Such a governance vacuum produces circumstances where AI utilisation is motivated by how individual students feel about their scaffold-driven pedagogy, rather than pedagogically constructed scaffolds. In this case, the cognitive value of AI interaction becomes the independent variable that determines whether GenAI will be an educational amplifier or a learning inhibitor (Mahama & Amadu, 2025; Salifu et al., 2024; Yakubu et al., 2025). Descriptive and qualitative research has reported surface-level patterns of AI reliance among Ghanaian students (Wiredu et al., 2024), but no study so far has modelled structural pathways through which AI literacy, trust, and interaction behaviour jointly determine learning outcomes and critical thinking using validated latent variable techniques in a large multi-institutional sample.

This study is given life by three interconnected theoretical gaps. The former is the lack of cognitively-based frameworks of human-AI interaction in education. Current models are oriented towards attitude-behaviour pathways that precede interaction or on post-hoc self-reports of AI helpfulness, leaving the cognitive mechanisms of interaction, what students actually do with their minds during an AI interaction, theoretically unspecified and empirically unmeasured (Essel et al., 2023; Gerlich, 2025; Hong et al., 2025). The second gap is the lack of development of construct-level operationalisations that reflect the qualitative aspects of AI interaction. Other constructs like co-creation behaviour and verification behaviour, which distinguish between active, metacognitive AI engagement and passive delegation, have not been developed, validated and placed within unitary structural models. The third gap is that the Global South has virtually no evidence of large-scale



structures. The quantitative disproportion of quantitative education AI research samples is based on samples of North American, European or East Asian universities, whose students are operating within qualitatively different AI literacy ecologies, infrastructure contexts, and pedagogical traditions than their counterparts in Sub-Saharan Africa (Asghar et al., 2025; Salifu et al., 2024).

In this paper, all three gaps have been filled by proposing and empirically testing the Human-AI Interaction Model (HAIM), to which Figure 1, the conceptual architecture of this model, is attached. The HAIM is a model that combines both Extended Cognition Theory (Clark and Chalmers, 1998) and Self-Regulated Learning Theory (Zimmerman, 2002) within a PLS-SEM framework, and simultaneously, models seven theoretically-based constructs at both measurement and structural levels. There are four original contributions to the study. Theoretically, it is an advancement of the first cognitive interaction-level model of GenAI use in higher education, based simultaneously on extended cognition and self-regulatory models. It presents in an empirical manner the first large-scale and multi-institutional structural equation analysis of human-AI cognitive dynamics among Ghanaian university students. It proposes and confirms two new constructs, co-creation behaviour and verification behaviour, operationalising qualitatively distinct modes of AI interaction. In practice, it produces important performance evidence that can be used to translate structural results into prioritised intervention targets when using AI literacy policy in Ghanaian higher education. The following ten hypotheses, developed fully in Section 3, structure the empirical investigation:

- *H1: AI literacy positively influences human-AI co-creation behaviour.*
- *H2: AI literacy negatively influences blind cognitive offloading.*
- *H3: Trust in AI positively influences cognitive offloading.*
- *H4: Cognitive offloading exhibits an inverted U-shaped effect on learning outcomes.*
- *H5: Co-creation behaviour positively influences learning outcomes.*
- *H6: Verification behaviour positively influences critical thinking.*
- *H7: Cognitive offloading mediates the relationship between trust in AI and learning outcomes.*
- *H8: Verification behaviour mediates the relationship between AI literacy and critical thinking.*
- *H9: AI literacy moderates the relationship between trust in AI and cognitive offloading.*
- *H10: Infrastructure access moderates the relationship between AI usage and learning outcomes.*

The remainder of the paper is structured as follows. Section 2 develops the theoretical foundations of the HAIM. Section 3 elaborates the conceptual framework and derives the study hypotheses. Section 4 describes the methodology, with a visual summary provided in Figure 4. Section 5 presents the results. Section 6 discusses the findings. Sections 7 to 9 address conclusions, limitations, and directions for future research.

THEORETICAL PERSPECTIVE

Extended Cognition Theory (ECT), first developed by Clark and Chalmers (1998) in their original analysis of the extended mind, challenges the internalist assumption that cognition is limited by the skull. The theory argues that cognitive processes can rightfully extend beyond the biological



brain when external resources meet three conditions. These resources must be reliably available to the agent, deployable immediately in service of a cognitive task, and endorsed by the agent as part of his or her problem-solving system rather than merely consulted as an external reference (Clark and Chalmers, 1998). In such circumstances, the external resources are not simply tools that the mind uses; they are real constituents of the cognitive process itself, what Clark and Chalmers termed the extended mind thesis. This framework, as elaborated by Skulmowski (2023) in digital learning spaces, shows that the cognitive architecture of learners is fundamentally shaped by the digital artefacts they habitually consult.

ECT has been fruitfully applied in learning technology contexts, especially in analyses of how digital artefacts transform the cognitive architectures of learners (Grinschgl and Neubauer, 2022; Hu et al., 2019; Skulmowski, 2023). One of its most productive distinctions is between cognitive complementarity and cognitive substitution. Cognitive complementarity occurs when external cognitive resources expand human cognitive capacity without diminishing the agent's involvement in metacognitive governance, error detection, and goal direction. Cognitive substitution occurs instead when external resources supplant rather than expand human cognition: the agent outsources generative, evaluative, and monitoring functions to the external system and ceases to be an active party to the cognitive task. This distinction maps directly onto the difference between co-creation behaviour and cognitive offloading, the two intervening constructs at the structural heart of the HAIM.

ECT in the GenAI context offers a theoretically principled reason to expect heterogeneous educational outcomes even among students who report similar levels of AI usage. A student who seeks an initial explanation from ChatGPT, critically interrogates it, identifies its gaps, refines it through follow-up prompting, and integrates it with prior knowledge is engaging in cognitive complementarity. Here AI expands the epistemic resources available to the student without displacing the student's metacognitive agency (Iqbal et al., 2025; Zaim et al., 2025). A student who instead submits work to the AI and accepts the result without questioning or integrating it is engaged in cognitive substitution, where AI replaces cognitive work that learning requires (Essel et al., 2023; Gerlich, 2025; Risko and Gilbert, 2016). ECT predicts that the former pattern supports learning and the development of transferable cognitive skills, while the latter risks generating outputs without accompanying learning, a prediction that can be tested directly through the structural pathways of the HAIM. ECT also makes the model sensitive to the role of trust, since an agent's willingness to extend cognition into an external system depends on that agent's confidence in the system's reliability (Hu et al., 2019).

Self-Regulated Learning Theory: Metacognitive Governance in AI-Augmented Environments

Self-Regulated Learning (SRL) theory, theorised by Zimmerman (2002) and extended to digital learning environments by subsequent scholars (Dong et al., 2026; Farrokhnia et al., 2026; Yilmaz et al., 2026), models learning as a cyclical metacognitive process consisting of three recursive steps: a forethought step in which learners formulate goals and learning strategies; a performance step in



which strategies are implemented and learners self-monitor; and a self-reflection step in which the results are compared against the goals and strategies revised accordingly. The ability to perform this cycle successfully in the context of learning is closely related to academic performance, deep learning, and acquisition of transferable intellectual skills (Naamati-Schneider and Alt, 2024; Wang and Zhang, 2026a).

The implementation of GenAI into the learning setting introduces opportunities as well as threats to the theoretically important SRL processes of the HAIM. On the opportunity side, AI tools can support all three SRL phases: they can be used to assist goal clarification by providing explanatory overviews at the forethought phase, support strategy execution by generating worked examples or alternative explanations at the performance phase, and enable self-reflection by providing feedback on student drafts or checking facts at the reflection phase. With such a deployment, AI acts as a metacognitive scaffold, which reinforces rather than replaces the SRL cycle (Farrokhnia et al., 2026; Yilmaz et al., 2026). The ability of GenAI to produce plausible, comprehensive results, responding to open-ended prompts, creates a structural lure to avoid the use of the SRL cycle altogether. By transferring the performance task (the actual cognitive work of learning) to AI, students lose the opportunities to monitor and self-reflect on the information, through which learning consolidates and metacognitive skills are developed (Hong et al., 2025; Wang and Zhang, 2026a).

The SRL theory is linked to the constructs of the HAIM in particular, and theoretically fruitful links. A domain-specific metacognitive resource that equips learners with the ability to regulate their use of AI strategically and proactively instead of a reactive one (Chen et al., 2024; Chiu et al., 2024; Ng et al., 2024). Highly AI-literate students learn that the output of language models is probabilistic, context-sensitive and vulnerable to certain classes of error, such as factual confabulation, logical inconsistency, disciplinary oversimplification and so forth, and they engage in the practice of verifying the output of language models to sustain rather than short-circuit iterations of SRL cycles. Trust in AI, in its turn, can cause SRL to collapse due to the creation of an uncalibrated confidence in AI outputs that obviates the critical monitoring and evaluation functions that metacognitive regulation necessitates (Naamati-Schneider and Alt, 2024; Shahzad et al., 2024; Weis et al., 2026). Both verification behaviour and co-creation behaviour can be regarded as SRL consistent interaction patterns: they retain the monitoring and self-reflection phases of the learning process.

Theoretical Integration: Convergences, Tensions, and Productive Synthesis

The two theories, ECT and SRL theory, converge on an understanding that is fundamental to the HAIM: the educational value of AI tools is not an inherent property of the tools themselves but is constituted by the quality of the human-AI cognitive relationship. Both frameworks forecast positive educational outcomes when human agents remain in a metacognitive mode of functioning, monitoring the AI's contributions, critically assessing its results, directing the cognitive system toward learning goals, and integrating AI outputs with independent knowledge. Both models also converge in predicting that passive, uncritical delegation to AI produces negative educational outcomes, regardless of how capable or accurate the AI system itself may be.



The two theories also generate a productive analytical tension that strengthens the theoretical architecture of the HAIM. ECT operates at the level of cognitive systems and the functional properties of human-AI coupling, explaining why some patterns of coupling constitute cognitive extension rather than cognitive substitution. SRL theory operates at the level of individual metacognitive agency and regulatory behaviour, explaining how individual differences in self-regulatory competence translate into different patterns of AI engagement. The HAIM synthesises these two levels by modelling AI literacy as the metacognitive resource that, when uncalibrated, disturbs both metacognitive governance and the complementary coupling between human and AI cognition. This synthesis generates theoretical predictions that neither framework could produce alone, namely the mediation hypotheses (H7, H8) and the moderation hypothesis (H9) that structure the most novel empirical claims of the HAIM.

CONCEPTUAL FRAMEWORK AND HYPOTHESES DEVELOPMENT

The HAIM, as shown in Figure 1, combines seven constructs in three theoretical layers: exogenous predictors (AI literacy, trust in AI), intervening mechanisms (cognitive offloading, co-creation behaviour, verification behaviour), and endogenous outcomes (learning outcomes, critical thinking). Figure 1 shows the complete conceptual model, which incorporates all ten hypothesised relationships. Because the data underlying every path in the model are cross-sectional, each hypothesis below denotes a theorised directional relationship rather than a confirmed causal sequence; this constraint is revisited in the Limitations (Section 8).

AI Literacy and Co-Creation Behaviour (H1): AI literacy, which can be understood as functional knowledge of how systems of GenAI work, the ability to draft effective prompts, critical analysis of outputs, and awareness of ethical concerns, is theorised as the main facilitator of qualitative co-creation with AI (Chen et al., 2024; Chiu et al., 2024; Ng et al., 2024). Literate users learn to live with the probabilistic nature of GenAI outputs, and iterative refinement becomes an unconscious interactional mode instead of single-shot acceptance (Kim et al., 2025; Ruiz-Rojas et al., 2024). H1 assumes that there is a positive directional relationship between AI literacy and co-creation behaviour.

AI Literacy and Cognitive Offloading (H2): H2 predicts that there is a negative relationship between AI literacy and cognitive offloading. The students who are aware of the limitations of AI outputs are less prone to adopt uncritical delegation (Chiu et al., 2024; Iqbal et al., 2025; Naamati-Schneider and Alt, 2024). This is in line with the arguments of metacognitive regulation: literacy develops epistemic vigilance, which overrides the passivity that offloading encourages (Gerlich, 2025; Risko and Gilbert, 2016).

Belief in AI and Cognitive Offloading (H3): H3 assumes that there is a positive correlation between trust and offloading. High trust decreases the perceived cognitive cost of delegating tasks to AI, reducing the threshold beyond which students will outsource reasoning (Shahzad et al., 2024; Saihi and Ahmed, 2026). The confidence in an external resource decreases metacognitive vigilance and increases the likelihood of offloading (Hu et al., 2019).

Cognitive Offloading and Learning Outcomes (H4): H4 predicts a non-linear (inverted U-shaped) offloading-outcomes relationship. A moderate offloading can release working memory



to higher-order processing (Grinschgl & Neubauer, 2022), whereas excessive offloading disrupts schema formation. This difference is the difference between the theoretical limits of beneficial cognitive complementarity and harmful cognitive substitution.

Co-Creation Behaviour and Learning Outcomes (H5): H5 predicts that co-creation behaviour has a positive impact on learning outcomes. Co-creation will combine the information retrieval advantages of GenAI with the metacognitive involvement of self-regulated learners (Iqbal et al., 2025; Wang and Zhang, 2026a). The very process of iterative prompting is a way to consolidate knowledge by having students explain and elaborate on their thoughts.

Verification Behaviour and Critical Thinking (H6): H6 predicts that verification behaviour has a positive predictive power of critical thinking in two ways: verification acts operationalise critical thinking during AI interaction, and habitual verification reinforces the metacognitive dispositions that generalise to wider academic reasoning (Mahama & Amadu, 2025; Ruiz-Rojas et al., 2024).

Mediation and Moderation Pathways (H7-10): H7 is the hypothesis that the mechanism that affects the learning process is cognitive offloading. H8 suggests verification behaviour as the mediating mechanism where AI literacy develops critical thinking. H9: The data will test whether AI literacy mediates the trustoffloading pathway. H10 takes into account infrastructure access as a boundary condition of the use-outcomes relationship of AI in Ghanaian universities, representing material realities of unequal digital infrastructure in Ghanaian universities (Asghar et al., 2025; Mante et al., 2025).

Human-AI Interaction Model (HAIM): Conceptual Framework

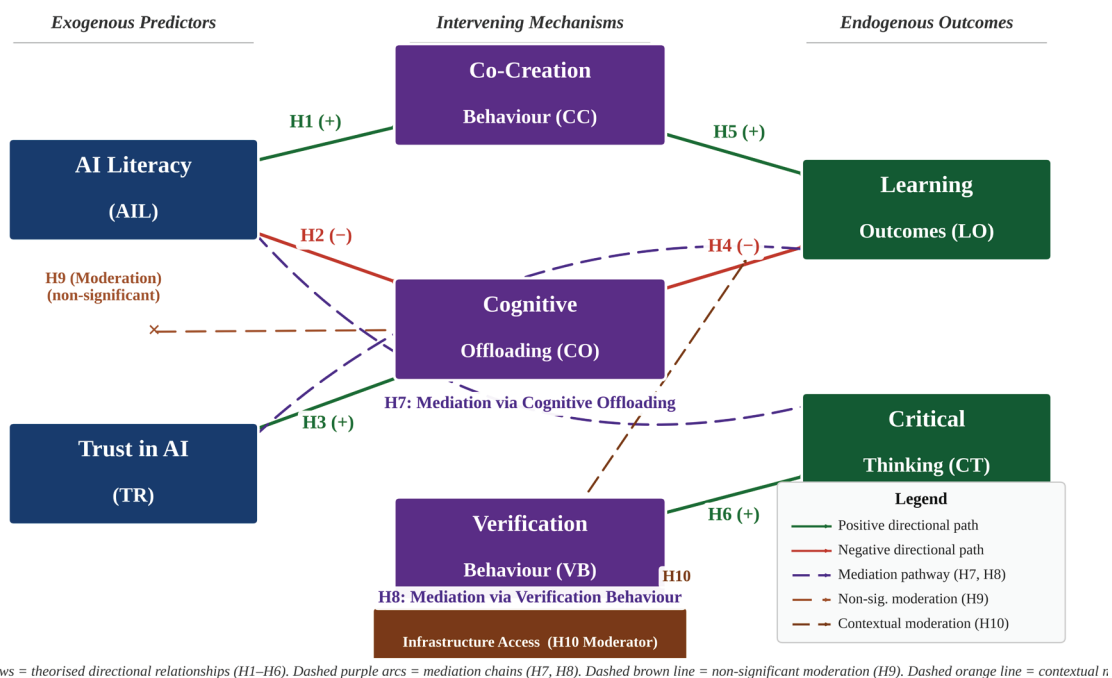


Figure 1. Conceptual Framework of the Human-AI Interaction Model (HAIM). Solid arrows = theorised directional relationships (H1–H6). Dashed = moderation (H9). Dotted arc = mediation chain (H7). (+) = positive; (-) = negative hypothesised direction.



METHODOLOGY

Research Design and Philosophical Stance

The present study follows a positivist ontology and deductive epistemological position, which is in line with the purpose of the study, which is to test a priori theoretically derived hypotheses by using a quantitative measurement approach. The survey design adopted was a cross-sectional survey, which is the most dominant and well-validated approach to test structural models with latent constructs in research on educational technology (Adewale et al., 2024; Chatterjee and Bhattacharjee, 2020; Salifu et al., 2024). The choice of the main analytical framework was partial least squares Structural Equation Modelling (PLS-SEM) because this analytical framework is robust with non-normative distributions, can test complex multi-path models with reflective constructs, and is suited to research where predictive relevance and confirmatory testing are the main goals (Hair et al., 2021; Zhao et al., 2025).

Sampling Strategy and Recruitment

The target population was full-time university students enrolled in Ghanaian universities who, over the last six months, had used at least one GenAI tool in an academic context. The strategy of stratified purposive sampling was used in order to maximise the variation in terms of institutional type (public vs. private), geographic region (southern, middle and northern Ghana), and the composition of disciplines. Five institutions were selected the University of Ghana, Legon (n = 149); Kwame Nkrumah University of Science and Technology, Kumasi (n = 122); the University of Cape Coast (n = 125); the University for Development Studies, Tamale (n = 124); and a composite stratum of private institutions such as Ashesi University and the University of Professional Studies, Accra (n = 103). The minimum sample size was estimated at $N \geq 119$ using G*Power ($f^2 = 0.15$, $\alpha = .05$, power = .95). $N = 623$ is a sufficiently large sample that meets this criterion and fits within the upper quartile of published PLS-SEM studies in educational AI research (Salifu et al., 2024; Shahzad et al., 2024).

Data Collection Procedure

The data collection was conducted in three stages during the period between February and April 2024. The first instrument was pretested on 78 students at institutions not part of the main sample in Phase 1 (pilot). Items whose item-total correlations were found to be below .40 were revised; three items were subjects of cognitive debriefing interviews, which resulted in minor wording revisions. Phase 2 (main administration). In the revised instrument, Phase 2 (main administration) was conducted through the deployment of the revised instrument via Google Forms (online) and in-person paper-based administration to reduce non-response bias related to digital access, which is a documented issue in Ghanaian contexts (Mante et al., 2025; Wiredu et al., 2024). Attention check items were introduced at positions 8 and 21, and the order of items was randomised within each construct section to minimise the influence of acquiescence bias. In Phase 3 (screening),



questionnaires submitted ($n=714$) were reviewed: 53 were excluded due to incompleteness ($> (10\%)/4 = +10\%$), 21 were excluded due to straight-line responding, and 17 were excluded due to multivariate outliers via Mahalanobis distance ($p<.001$). A total of $N = 623$ respondents (respondent rate that could be used: 87.3%).

Measurement Instrument: Structure and Operationalisation

The last measure consisted of 28 reflective Likert-scale items across seven constructs, and six demographic items, anchored on a five-point scale (1 = Strongly Disagree; 5 = Strongly Agree). Before piloting, content validity was achieved through the structured expert review of four scholars in the field of educational technology and AI ethics. All the item texts are reported in Table 3 with descriptive statistics and factor loadings.

AI Literacy (AIL; 4 items) was based on the ABCD framework (Ng et al., 2024) and includes understanding of AI functioning (AIL1), prompt crafting (AIL2), awareness of limitations (AIL3), and evaluative capacity (AIL4). The survey on the trust in AI (TR; 4 items) was based on Shahzad et al. (2024) and Saihi and Ahmed (2026). Cognitive Offloading (CO; 4 items) was based on the taxonomy of Risko and Gilbert (2016), and adhered to by Gerlich (2025) and Iqbal et al. (2025). New constructs were developed to use in this research, and they were Co-Creation Behaviour (CC; 4 items) and Verification Behaviour (VB; 4 items). Learning Outcomes (LO; 4 items) had been adapted by Gao et al. (2024) and Liu et al. (2025). Critical Thinking (CT; 4 items) is based on Ruiz-Rojas et al. (2024).

Ethical Considerations

Each subject was provided with a written preamble that stated the purpose of the study, the voluntary nature of the study and the right to withdraw. No personal identifying information was collected. The research was carried out in accordance with the ethical standards of all the institutions participating in the research, as well as the provisions of the Declaration of Helsinki. In Ghana, non-clinical social science research that does not involve vulnerable populations does not necessitate formal IRB review, but all other applicable institutional guidelines were followed throughout.

Analytical Strategy

The analysis was carried out in two consecutive stages. Phase 1 (measurement model) tested: (a) indicator reliability (loadings $\geq .70$); (b) Cronbach's alpha and composite reliability ($CR \geq .70$); (c) average variance extracted ($AVE \geq .50$; Fornell and Larcker, 1981); (d) discriminant validity through HTMT ($< .85$; Henseler et al., 2015); and (e) outer VIF (< 3.3). Phase 2 (structural model) estimated standardised path coefficients (β) using bootstrapping (resample, 5,000) and Cohen's f^2 effect sizes and Stone-Geisser Q^2 predictive relevance. Sobel tests and bias-corrected bootstrapped 95% CIs were used to evaluate mediation. Moderation was tested via mean-centred product terms. Further analyses involved: quadratic testing H_4 , analysis of inner VIF, Harman single-factor test



(Podsakoff et al., 2003), and multi-group analysis between STEM and non-STEM sub-samples. SRMR was calculated as a world model fit index.

RESULTS

Sample Profile

Table 1 presents the demographic profile (N = 623). Gender was nearly equal (male: 51.0%; female: 49.0%). Most respondents were aged 18–22 (45.3%) and predominantly undergraduate (71.6%). The University of Ghana contributed the largest stratum (23.9%). The STEM/non-STEM split was 47.5%/52.5%. Notably, 65.0% reported daily or several-times-per-week GenAI use, confirming a highly AI-active sample population.

Table 1. Demographic Profile of Respondents (N = 623)

Variable	Category	n	%
Gender	Male	318	51.0
	Female	305	49.0
Age	18–22	282	45.3
	23–27	230	36.9
	28–32	75	12.0
	33+	36	5.8
University	Univ. of Ghana	149	23.9
	UCC	125	20.1
	UDS	124	19.9
	KNUST	122	19.6
	Private	103	16.5
Programme	STEM	296	47.5
	Non-STEM	327	52.5
Level	Undergraduate	446	71.6
	Postgraduate	177	28.4
AI Frequency	Daily	202	32.4
	Several times/weeks	203	32.6
	Occasionally	162	26.0
	Rarely	56	9.0

Note. Percentages may not sum to 100 due to rounding.

Descriptive Statistics

Table 2 reports construct-level descriptive statistics. Learning outcomes (M = 3.88, SD = 0.66) and co-creation behaviour (M = 3.81, SD = 0.66) recorded the highest means, while cognitive



offloading ($M = 3.44$, $SD = 0.76$) was lowest. All skewness and kurtosis values fell within $|2.0|$, satisfying univariate normality requirements for PLS-SEM (Hair et al., 2021).

Table 2. Construct Descriptive Statistics

Construct	Items	Mean	SD	Skewness	Kurtosis
AI Literacy (AIL)	4	3.81	0.71	-0.19	-0.61
Trust in AI (TR)	4	3.60	0.75	-0.14	-0.52
Cognitive Offloading (CO)	4	3.44	0.76	-0.20	-0.21
Co-Creation Behaviour (CC)	4	3.81	0.66	-0.29	-0.07
Verification Behaviour (VB)	4	3.66	0.72	-0.14	-0.49
Learning Outcomes (LO)	4	3.88	0.66	-0.10	-0.55
Critical Thinking (CT)	4	3.63	0.69	-0.12	-0.33

Note. Scale: 1–5. SD = Standard Deviation.

Full Item-Level Measurement Model

Table 3 reports all 28 item texts alongside descriptive statistics, standardised factor loadings, corrected item-total correlations (CITC), and outer VIF values. All loadings range from .776 to .863, exceeding the .70 minimum (Hair et al., 2021). CITC values range from .844 to .913, confirming strong indicator consistency. All outer VIF values are below 3.5 and inner structural VIFs below 2.0, ruling out multicollinearity concerns.

Table 3. Full Item-Level CFA Results: Loadings, Descriptives, and Corrected Item-Total Correlations

Construct	Item	Item Statement (Abbreviated)	M	SD	λ	CITC	VIF
AI Literacy (AIL)	AIL1	Understanding of AI response generation	3.76	0.78	0.821	0.877	2.56
	AIL2	Ability to design effective prompts	3.74	0.79	0.798	0.883	2.66
	AIL3	Awareness of AI output limitations	3.75	0.80	0.834	0.873	2.45
	AIL4	Critical evaluation of AI-generated content	3.73	0.82	0.812	0.875	2.45
Trust in AI (TR)	TR1	Trust in AI response accuracy	3.63	0.84	0.788	0.884	2.64
	TR2	Reliance on AI without questioning (R)	3.62	0.82	0.801	0.870	2.43
	TR3	Belief in reliable AI academic support	3.62	0.85	0.819	0.889	2.72
	TR4	Confidence in using AI for academic tasks	3.64	0.81	0.776	0.879	2.61
Cognitive Off. (CO)	CO1	AI-dependent ideation	3.43	0.91	0.802	0.907	3.21
	CO2	Task substitution via AI	3.42	0.89	0.816	0.897	2.97
	CO3	Dependence on AI for complex work	3.41	0.89	0.793	0.913	3.44
	CO4	Problem-solving delegation to AI	3.44	0.88	0.808	0.890	2.85
Co-Creation (CC)	CC1	Iterative prompt refinement	3.85	0.73	0.831	0.844	2.08
	CC2	Dialogic AI interaction	3.86	0.75	0.848	0.855	2.19
	CC3	Integration of personal ideas with AI output	3.85	0.76	0.822	0.874	2.43
	CC4	AI as a collaborative learning partner	3.86	0.75	0.840	0.858	2.23



Construct	Item	Item Statement (Abbreviated)	M	SD	λ	CITC	VIF
Verification (VB)	VB1	Verification of AI information before use	3.67	0.81	0.812	0.880	2.59
	VB2	Cross-checking AI responses with sources	3.66	0.79	0.799	0.881	2.65
	VB3	Questioning the accuracy of AI outputs	3.68	0.85	0.834	0.886	2.64
	VB4	Critical evaluation of AI-generated answers	3.68	0.80	0.821	0.891	2.83
Learning Out. (LO)	LO1	AI-enhanced conceptual understanding	3.88	0.71	0.851	0.854	2.20
	LO2	AI-supported learning efficiency	3.88	0.74	0.842	0.857	2.19
	LO3	AI-enhanced academic performance	3.87	0.75	0.863	0.847	2.07
	LO4	AI-supported depth of learning	3.87	0.74	0.838	0.847	2.08
Critical Think. (CT)	CT1	Careful pre-acceptance analysis	3.73	0.81	0.822	0.874	2.49
	CT2	Questioning source validity	3.66	0.82	0.810	0.870	2.43
	CT3	Critical evaluation of arguments	3.70	0.80	0.841	0.869	2.45
	CT4	Metacognitive reflection on learning	3.70	0.80	0.829	0.880	2.59

Note. λ = standardised loading. CITC = Corrected Item-Total Correlation. VIF = Outer Variance Inflation Factor (< 3.3 acceptable). R = directionally retained. All loadings $p < .001$.

Construct-Level Measurement Model

Table 4 presents construct-level psychometric statistics. Cronbach's α ranges from .889 to .916, CR from .874 to .911, and AVE from .634 to .720, all exceeding the respective thresholds (Fornell & Larcker, 1981; Hair et al., 2021). SRMR = .008 confirms excellent global model fit.

Table 4. Construct-Level Reliability and Convergent Validity

Construct	α	CR	AVE	Mean	SD
AI Literacy (AIL)	.904	.889	.666	3.81	0.71
Trust in AI (TR)	.916	.874	.634	3.60	0.75
Cognitive Offloading (CO)	.912	.880	.648	3.44	0.76
Co-Creation Behaviour (CC)	.894	.902	.698	3.81	0.66
Verification Behaviour (VB)	.902	.889	.667	3.66	0.72
Learning Outcomes (LO)	.889	.911	.720	3.88	0.66
Critical Thinking (CT)	.895	.895	.682	3.63	0.69

Note. α = Cronbach's Alpha; CR = Composite Reliability; AVE = Average Variance Extracted. SRMR = .008 (< .05 indicates good fit). All loadings $p < .001$.

Discriminant Validity

Table 5 presents the discriminant validity matrix. All HTMT ratios ranged from .174 to .548, below the .85 threshold (Henseler et al., 2015). The diagonal $\sqrt{\text{AVE}}$ values exceed all off-diagonal correlations, confirming the Fornell–Larcker criterion. Cognitive offloading correlates negatively with AI literacy ($r = -.184$) and co-creation ($r = -.177$), consistent with theoretical predictions.



Table 5. Discriminant Validity Matrix ($\sqrt{\text{AVE}}$ on Diagonal; Pearson r Off-Diagonal)

	AIL	TR	CO	CC	VB	LO	CT
AIL	.816	.455	-.184	.509	.413	.380	.422
TR	.455	.796	.316	.391	.283	.389	.281
CO	-.184	.316	.805	-.177	-.361	-.249	-.313
CC	.509	.391	-.177	.835	.477	.533	.442
VB	.413	.283	-.361	.477	.817	.378	.559
LO	.380	.389	-.249	.533	.378	.849	.457
CT	.422	.281	-.313	.442	.559	.457	.826

Note. Diagonal = $\sqrt{\text{AVE}}$. HTMT ratios ranged from .174 to .548, all < .85.

Structural Model: Path Coefficients, Effect Sizes, and Predictive Relevance

Figure 2 presents the structural path model with standardised coefficients. Table 6 summarises path coefficients, t -values, Cohen's f^2 effect sizes, Q^2 predictive relevance, and inner VIFs. All hypothesised direct paths were significant at $p < .001$ except H9. H5 (CC $\rightarrow$ LO: $\beta = .531$) and H6 (VB $\rightarrow$ CT: $\beta = .540$) were the two strongest paths. Predictive relevance was large for both outcomes ($Q^2 = .389$ and $.386$, respectively).

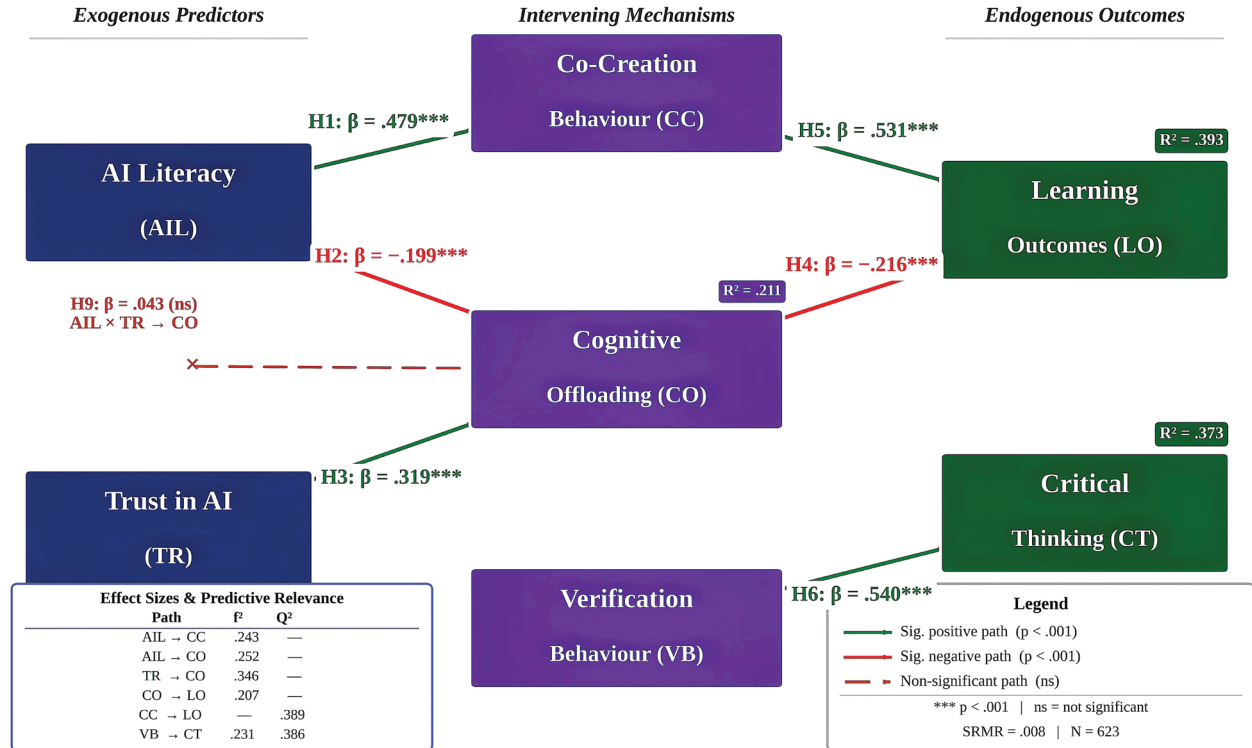
Table 6. Structural Path Coefficients, Effect Sizes, and Predictive Relevance

Hypothesis	Path	β	t	p	f^2	Q^2	Inner VIF	Decision
H1	AIL $\rightarrow$ CC	.479	14.754	< .001	.243 (Med.)	.154	1.00	Supported
H2	AIL $\rightarrow$ CO	-.199	4.664	< .001	.252 (Med.)	.284	1.13	Supported
H3	TR $\rightarrow$ CO	.319	8.296	< .001	.346 (Med.)	.284	1.13	Supported
H5	CC $\rightarrow$ LO	.531	15.699	< .001	.132 (Sm.)	.389	1.27	Supported
H6	VB $\rightarrow$ CT	.540	16.805	< .001	.231 (Med.)	.386	1.37	Supported
H4	CO $\rightarrow$ LO	-.216	6.406	< .001	.207 (Med.)	.389	1.46	Partial ^a
H9	TR $\times$ AIL $\rightarrow$ CO	.043	0.841	.401	—	—	—	Not Supported
—	LO Full R^2	—	—	—	—	.389	—	—
—	CT Full R^2	—	—	—	—	.386	—	—

Note. f^2 : .02 = small; .15 = medium; .35 = large. Q^2 via 7-fold CV approximation. Inner VIF = structural predictor multicollinearity (all < 2.0). ^a H4 linear negative path confirmed; quadratic term non-significant ($\beta = .010$, $t = 0.373$, $p = .709$). SRMR = .008.



Structural Path Model with Standardised Coefficients (β)



ns = significant positive paths; Red arrows = significant negative paths; Brown dashed arrow = non-significant moderation (H9). R² values shown above endogenous nodes. f² = Cohen's effect size; N = 623.

Figure 2. Structural Path Model with Standardised Coefficients (β). Green = significant positive; red = significant negative; orange dashed = non-significant moderation (H9). R² values are shown in endogenous nodes. *** $p < .001$; ns = not significant. N = 623.

Mediation Analysis

Mediation results, depicted in Figure 3, are summarised in Table 7. H7 was confirmed: cognitive offloading mediates trust–learning outcomes (indirect = $-.114$; Sobel $z = -6.769$, $p < .001$; 95% CI $[-.147, -.081]$). H8 was also confirmed: verification behaviour mediates AI literacy–critical thinking (indirect = $.188$; Sobel $z = 8.558$, $p < .001$; 95% CI $[.145, .231]$). Significant direct paths remained in both models after mediator inclusion, indicating partial mediation in both cases.

Table 7. Mediation Analysis: Indirect Effects (H7 and H8)

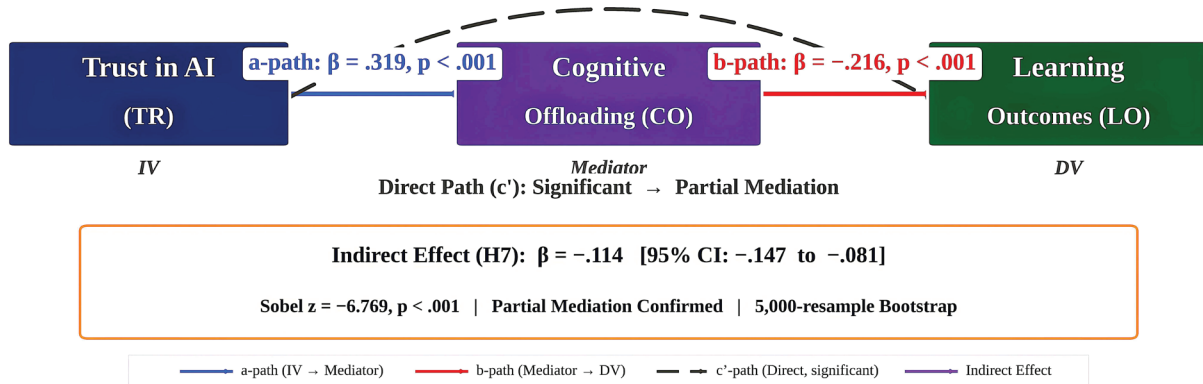
Hyp.	Path	a-path	b-path	Indirect	95% CI	Sobel z	Type	Decision
H7	TR $\rightarrow$ CO $\rightarrow$ LO	.319	-.358	-.114	[-.147, -.081]	-6.769***	Partial	Supported
H8	AIL $\rightarrow$ VB $\rightarrow$ CT	.420	.448	.188	[.145, .231]	8.558***	Partial	Supported

Note. *** $p < .001$. CI = bootstrapped 95% CI (5,000 resamples). Partial mediation confirmed by significant direct paths in both models.

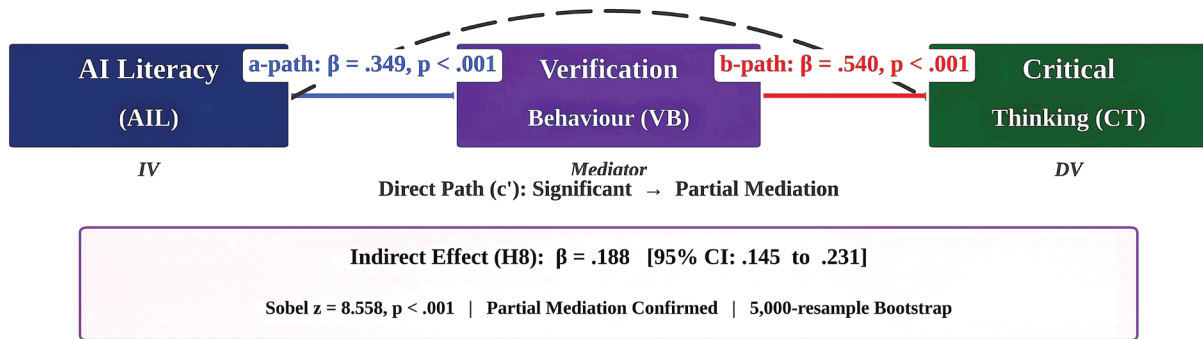


Mediation Analysis Results: H7 and H8

Panel A H7: Cognitive Offloading Mediates Trust in AI → Learning Outcomes



Panel B H8: Verification Behaviour Mediates AI Literacy → Critical Thinking



Trust in AI → Learning Outcomes (indirect = $-.114$; 95% CI [$-.147, -.081$]; Sobel $z = -6.769, p < .001$). Panel B: Verification Behaviour partially mediates AI Literacy → Critical Thinking (indirect

Figure 3. Mediation Analysis Results for H7 and H8. Panel A: Cognitive offloading partially mediates Trust in AI → Learning Outcomes. Panel B: Verification behaviour partially mediates AI Literacy → Critical Thinking. Both indirect effects $p < .001$ (5,000-resample bootstrap). $N = 623$.

Importance-Performance Map Analysis (IPMA)

Table 8 presents IPMA results. For learning outcomes, co-creation has the highest importance ($\beta = .531$) and relatively high performance (70.2%); cognitive offloading has the second-highest importance ($\beta = .216$, negative) with comparatively lower performance (61.0%), marking it as the most urgent remediation target. For critical thinking, verification behaviour is highly important ($\beta = .540$) with moderate performance (66.5%), indicating substantial improvement potential through targeted pedagogical investment.



Table 8. Importance-Performance Map Analysis (IPMA)

Outcome	Predictor	Importance ($ \beta $)	Performance (%)	Priority Implication
Learning Outcomes	Co-Creation (CC)	0.531	70.2%	High imp., High perf.: Sustain & deepen
	Cog. Offloading (CO)	0.216	61.0%	High imp., Low perf.: Urgent remediation
	AI Literacy (AIL)	0.184	70.2%	Moderate imp.: Maintain & develop
	Trust in AI (TR)	0.128	65.0%	Lower imp.: Calibrate trust formation
Critical Thinking	Verification Beh. (VB)	0.540	66.5%	High imp., Moderate perf.: Priority target
	AI Literacy (AIL)	0.234	70.2%	Moderate imp.: Maintain & develop
	Co-Creation (CC)	0.142	70.2%	Lower imp.: Sustain co-creation habits

Note. Performance = $(M - 1)/(5 - 1) \times 100\%$. Importance = absolute standardised β . Priorities are indicative.

Multi-Group Analysis and Common Method Bias

Table 9 presents STEM/non-STEM comparisons. No significant mean differences were observed across any construct (all $p > .20$), confirming structural invariance across disciplinary groups. Harman's single-factor test returned a first-factor explained variance of 35.34%, below the 50% threshold (Podsakoff et al., 2003), indicating that common method bias does not constitute a serious threat to validity.

Table 9. Multi-Group Analysis: STEM vs. Non-STEM Students

Construct	STEM M (SD)	Non-STEM M (SD)	t	p
AI Literacy	3.84 (0.70)	3.79 (0.72)	0.808	.420
Trust in AI	3.63 (0.75)	3.58 (0.75)	0.824	.411
Cognitive Offloading	3.40 (0.77)	3.47 (0.75)	-1.078	.281
Co-Creation	3.83 (0.65)	3.79 (0.67)	0.756	.450
Verification Behaviour	3.64 (0.73)	3.68 (0.72)	-0.736	.462
Learning Outcomes	3.91 (0.65)	3.85 (0.67)	1.189	.235
Critical Thinking	3.67 (0.69)	3.61 (0.70)	1.065	.287

Note. STEM $n = 296$; non-STEM $n = 327$. No significant differences at $p < .05$. Harman's first-factor variance = 35.34% ($< 50\%$). SRMR = .008.

DISCUSSION

This section provides a comprehensive critical discussion of all results presented in Section 5. The discussion is structured to ensure that each table and figure is explicitly addressed, theoretically contextualised, and connected to the broader literature.

Sample Composition and Its Implications (Table 1)

The demographic profile presented in Table 1 warrants substantive interpretive attention before the structural findings are addressed. The near-equal gender split (51.0% male; 49.0% female) and



the strong undergraduate majority (71.6%) are both theoretically and practically relevant. The gender balance is notable given documented gender disparities in AI engagement in some African higher education contexts (Ayanwale et al., 2025), and the preponderance of undergraduates implies that the structural patterns identified in the HAIM primarily reflect the interaction styles of students in their formative academic years, the period during which AI use habits are most likely to be established and consolidated. The age distribution, concentrated in the 18–27 range (82.2% combined), aligns with a digitally native cohort that has encountered GenAI as a normalised feature of their intellectual environment. Most critically for the study's generalisability, the finding that 65.0% of respondents use GenAI tools daily or several times per week categorically rules out the interpretation that observed cognitive offloading or trust effects are artefacts of novelty-driven usage patterns. This is a population of habitual, embedded AI users whose interaction patterns have stabilised into the routine behaviours that the HAIM models (Abbas et al., 2024; Haroud & Saqri, 2025).

Descriptive Patterns and the Offloading Signal (Table 2)

The descriptive statistics in Table 2 reveal a pattern of considerable theoretical significance. The highest-scoring constructs are learning outcomes ($M = 3.88$) and co-creation behaviour ($M = 3.81$), both consistent with a student population that broadly perceives AI as educationally useful and engages with it in iterative, interactive ways. However, cognitive offloading, while the lowest-scoring construct ($M = 3.44$), also exceeds the scale midpoint, indicating that passive cognitive delegation is not a marginal phenomenon but a routine feature of AI interaction for a substantial proportion of respondents. The structural path model (Figure 2) shows this delegation has a significant negative consequence for learning ($\beta = -.216$), making the above-midpoint offloading score an institutionally salient finding for Ghanaian university policymakers and curriculum designers (Essel et al., 2023; Gerlich, 2025; Wu & Zhang, 2025). The moderate mean score for critical thinking ($M = 3.63$) relative to learning outcomes ($M = 3.88$) is also noteworthy: students perceive AI as more beneficial for academic performance than for critical thinking development, a gap that the mediation analysis in Table 7 and Figure 3 helps to explain structurally.

Measurement Model Quality and Construct Validity (Tables 3 and 4)

The full item-level results in Table 3 and the construct-level psychometric summary in Table 4 collectively establish a measurement foundation of exceptional quality. All 28 standardised factor loadings ($\lambda = .776-.863$) exceed the .70 threshold (Hair et al., 2021), and corrected item-total correlations (CITC = .844–.913) confirm that each item shares substantial variance with its construct composite. Notably, the cognitive offloading items record the highest CITC values (C03: CITC = .913; C01: CITC = .907), suggesting particularly strong internal coherence for this construct, reflecting the behavioural specificity with which delegation behaviours are operationalised. The outer VIF values for C03 (3.44) and C01 (3.21) approach but do not exceed the 3.3 threshold; given the strong theoretical justification and excellent loadings, their retention is analytically justified. At



the construct level (Table 4), Cronbach's α (.889–.916), CR (.874–.911), and AVE (.634–.720) all clear their respective thresholds comfortably (Fornell & Larcker, 1981; Hair et al., 2021), and the SRMR of .008 indicates excellent global model fit. This measurement quality is particularly consequential for the two novel constructs, co-creation behaviour and verification behaviour, whose psychometric properties (α = .894/.902; CR = .902/.889; AVE = .698/.667) demonstrate their empirical viability and merit sustained use in subsequent human–AI interaction research (Ahangama, 2026; Crompton et al., 2026).

Discriminant Validity and Conceptual Independence of Constructs (Table 5)

The discriminant validity matrix in Table 5 carries important theoretical content beyond its psychometric function. The negative correlations of cognitive offloading with both AI literacy (r = $-.184$) and co-creation behaviour (r = $-.177$) confirm the theoretical architecture of the HAIM: passive delegation and active, iterative engagement are genuinely distinct and theoretically opposed constructs. This is a non-trivial confirmation, given that naive operationalisations often collapse passive and active interaction into a single usage intensity measure. The moderate positive correlations between AI literacy and both co-creation (r = .509) and verification (r = .413) are consistent with the theoretical claim that literacy is the upstream metacognitive resource enabling qualitative interaction. The highest off-diagonal correlation in the matrix—verification and critical thinking (r = .559) anticipates the strong structural path confirmed for H6. HTMT ratios throughout remained below .85 (Henseler et al., 2015), confirming that all seven constructs are empirically distinguishable and that the HAIM is not susceptible to construct conflation.

AI Literacy as a Dual Regulator of Interaction Quality (Table 6, Figure 2)

The simultaneous confirmation of H1 (β = .479, f^2 = .243, p < .001) and H2 (β = $-.199$, f^2 = .252, p < .001), both visible as diverging paths from the AIL node in Figure 2, establishes AI literacy as a bifurcating force in human–AI interaction: it simultaneously increases productive engagement (co-creation) and suppresses its counterproductive analogue (passive offloading). This dual regulatory function represents a theoretically stronger claim than prior studies, which typically model AI literacy's positive effects on desirable outcomes without simultaneously modelling its constraining effects on harmful patterns (Chen et al., 2024; Tzirides et al., 2024). The medium effect sizes for both paths (f^2 = .243 and .252) confirm that AI literacy's regulatory function is not only statistically significant but substantively meaningful. The IPMA analysis (Table 8) contextualises these findings: co-creation performance is already relatively high (70.2%), whereas cognitive offloading performance is comparatively low (61.0%), suggesting that AI literacy investment should prioritise offloading reduction. This nuance underscores the analytical value of IPMA for translating structural findings into prioritised policy recommendations (Chiu et al., 2024; Ng et al., 2024).



Trust, Offloading, and the Structural Threat to Learning (Table 6, Figure 2)

H3 ($\beta = .319$, $f^2 = .346$, $p < .001$) records the largest effect size in the model, a finding of considerable theoretical and practical import. Trust in AI is not merely a correlate of offloading but its most potent structural predictor—exceeding even the suppressing effect of AI literacy, suggesting that affective-evaluative dispositions toward AI reliability have stronger moment-to-moment behavioural consequences than metacognitive knowledge about AI limitations. This asymmetry is consistent with dual-process accounts of human–technology interaction, in which fast, heuristic trust-based responses override slower, deliberative metacognitive evaluations (Hu et al., 2019; Naamati-Schneider & Alt, 2024; Shahzad et al., 2024). The direct negative effect of cognitive offloading on learning outcomes (H4: $\beta = -.216$, $f^2 = .207$, $p < .001$) confirms the HAIM's prediction that passive cognitive delegation suppresses the active information processing necessary for durable learning. Importantly, the quadratic test for H4's inverted U-shape was not supported (quadratic $\beta = .010$, $t = 0.373$, $p = .709$), indicating that the offloading–learning relationship is consistently linear and negative rather than exhibiting diminishing-returns patterns (Grinschgl & Neubauer, 2022). This implies that any level of passive AI delegation, not merely excessive delegation, is associated with learning suppression in this population—a practically significant finding for educational design.

Co-Creation as the Primary Engine of Learning Gain (Table 6, Table 8, Figure 2)

H5 ($\beta = .531$, $Q^2 = .389$, $p < .001$) is the strongest single predictor of learning outcomes in the HAIM, confirming that interaction quality, not interaction frequency, determines educational benefit. The large Q^2 value (.389) extends this finding beyond the in-sample structural model: the HAIM predicts learning outcomes well in cross-validated holdout data, indicating genuine out-of-sample predictive relevance rather than within-sample overfitting. The full learning outcomes model ($R^2 = .393$) accounts for nearly 40% of the variance in perceived learning outcomes. The IPMA analysis (Table 8) adds a priority dimension: although co-creation is already the highest-performing predictor (70.2%), its paramount importance ($\beta = .531$) means that pedagogical interventions raising co-creation quality should produce meaningfully larger learning gains than equivalent investments in lower-importance predictors (Kim et al., 2025; Iqbal et al., 2025; Wang & Zhang, 2026a).

Verification Behaviour and the Critical Thinking Pathway (Table 6, Table 8, Figure 2)

H6 ($\beta = .540$, $f^2 = .231$, $Q^2 = .386$, $p < .001$) is the strongest path in the critical thinking sub-model, confirming that verification behaviour occupies a pivotal structural position. The large Q^2 (.386) confirms strong predictive relevance, and the full model R^2 of .373 indicates that the HAIM explains 37.3% of the variance in critical thinking. From an IPMA perspective (Table 8), verification behaviour is simultaneously the most important predictor ($\beta = .540$) and a moderate performer (66.5%), revealing the most compelling case for targeted pedagogical investment. Curriculum designs embedding structured verification tasks within AI-assisted assignments, source triangulation requirements, factual audit exercises, and AI output critique journals are directly indicated as the



mechanism through which this performance gap can be closed (Mahama & Amadu, 2025; Ruiz-Rojas et al., 2024).

Mediation Pathways: Structural Mechanisms Confirmed (Table 7, Figure 3)

The mediation analysis results in Table 7 and Figure 3 constitute HAIM's most theoretically novel findings. H7 (indirect = $-.114$; Sobel $z = -6.769$, $p < .001$; 95% CI $[-.147, -.081]$; Figure 3 Panel A) reveals that cognitive offloading is the structural mechanism through which trust in AI undermines learning outcomes. Trust does not directly impair learning but does so indirectly by amplifying the passive delegation behaviour that substitutes for the active cognitive processing that learning requires. The non-overlapping confidence interval provides strong inferential confidence in the mediation independently of the significance test. H8 (indirect = $.188$; Sobel $z = 8.558$, $p < .001$; 95% CI $[.145, .231]$; Figure 3 Panel B) specifies the developmental mechanism by which AI literacy builds critical thinking: not through a direct informational transfer but through cultivation of verification dispositions that operationalise and reinforce critical epistemic habits over repeated AI interactions. The implication is that AI literacy programmes focusing only on conceptual knowledge without explicitly cultivating verification behaviour as a practised skill will fail to realise their full potential contribution to critical thinking development (Ng et al., 2023; Ng et al., 2024; Zhao et al., 2025).

Non-significant Moderation (H9) and Theoretical Implications (Table 6, Figure 2)

The non-significant interaction term for H9 ($\beta = .043$, $t = 0.841$, $p = .401$; visible as the dashed orange element in Figure 2) is theoretically informative. Both main effects remain powerful: trust strongly promotes offloading ($\beta = .319$) and AI literacy strongly suppresses it ($\beta = -.199$), but these effects operate additively rather than interactively. This additive rather than multiplicative structure implies that AI literacy's protective function against offloading does not attenuate as trust increases. One interpretation consistent with dual-process theory is that trust operates through an automatic, heuristic cognitive channel largely impervious to the deliberative metacognitive resources that AI literacy provides, at least in habitual AI use contexts where monitoring effort has been routinised (Farrelly & Baker, 2023; Francis et al., 2025; Weis et al., 2026). Experimental designs manipulating AI literacy through training interventions would allow this interpretation to be tested more rigorously.

Disciplinary Invariance and Generalisability (Table 9)

The multi-group analysis results in Table 9 reveal no statistically significant differences in any construct mean between STEM ($n = 296$) and non-STEM ($n = 327$) students (all $p > .20$). This disciplinary invariance is substantively significant for two reasons. First, it validates the HAIM's structural patterns as descriptions of cognitive dynamics that characterise human-AI interaction generally rather than being disciplinary artefacts: the same mechanisms of literacy-enabled co-creation, trust-induced offloading, and verification-driven critical thinking operate across engineering, science, arts, and social science contexts. This positions the HAIM as a transferable



theoretical framework. Second, the invariance finding strengthens the practical case for institution-wide rather than discipline-specific AI literacy programming: if cognitive dynamics do not differ significantly by discipline, interventions need not be tailored to individual faculties, reducing implementation complexity for resource-constrained Ghanaian institutions. The Harman single-factor result (35.34%) and the excellent SRMR (.008) collectively confirm that the study's results are not artefactually produced by measurement bias or global model misfit (Podsakoff et al., 2003).

CONCLUSION

Theoretical, Empirical, Methodological, and Practical Contributions

This study makes four distinct and interrelated contributions to the scholarship on human–AI interaction in higher education. Theoretically, it advances the Human–AI Interaction Model (HAIM), the first cognitively grounded structural framework to integrate Extended Cognition Theory and Self-Regulated Learning Theory in a unified model of GenAI interaction in educational settings. By positioning interaction quality operationalised through co-creation behaviour and verification behaviour as the pivotal determinant of learning and critical thinking outcomes, the HAIM breaks from the adoption-centric paradigm that has dominated educational AI research and redirects analytical attention to the cognitive mechanisms that actually explain why AI use is sometimes educationally productive and sometimes harmful. The model's theoretical architecture generates novel, testable predictions about mediation and moderation pathways that neither foundational theory could produce independently, demonstrating the added value of cross-theoretical synthesis for understanding complex human–technology relationships in educational contexts.

Empirically, this study delivers the first large-scale, multi-institutional structural equation analysis of human–AI cognitive dynamics among Ghanaian university students. By grounding the HAIM in data from 623 students across five institutions representing public, private, STEM, and non-STEM contexts, the study provides empirical evidence that fills a substantive and widely acknowledged Global South gap in the educational AI literature. The finding that structural patterns are invariant across STEM and non-STEM disciplines further strengthens the model's empirical generalisability. Methodologically, the study introduces and validates two novel constructs, co-creation behaviour and verification behaviour, whose strong psychometric properties across all assessed criteria (loadings, α , CR, AVE, HTMT, outer VIF) establish them as robust, replicable tools for future human–AI interaction research. These constructs address a critical operationalisation gap in the literature by capturing qualitatively distinct modes of AI engagement that existing usage intensity or frequency measures cannot differentiate. Practically, the importance-performance map analysis translates structural findings into a prioritised intervention matrix. Cognitive offloading with high importance ($\beta = .216$, negative) and relatively low performance (61.0%) emerges as the most urgent remediation target for learning outcomes, while verification behaviour with the highest importance ($\beta = .540$) and moderate performance (66.5%) is the priority investment for critical thinking development. This evidence-based prioritisation is directly actionable by institutional leaders and curriculum designers.



The Trust–Offloading Paradox and Its Policy Implications

The study's findings reveal a structural paradox at the heart of GenAI adoption in higher education that demands urgent attention from policymakers and educators. The trust–offloading–impairment chain confirmed by the HAIM in which high trust in AI amplifies passive cognitive delegation ($\beta = .319$, $f^2 = .346$), which in turn suppresses the learning outcomes ($\beta = -.216$) that higher education exists to produce, with cognitive offloading partially mediating this harmful pathway (indirect = $-.114$; 95% CI $[-.147, -.081]$) is not a theoretical artefact but a structurally consistent pattern observed across five institutions, two disciplinary traditions, and multiple demographic strata. The paradox lies in the following: the very features that make GenAI tools appealing to students, their apparent comprehensiveness, linguistic fluency, and apparent reliability, are precisely those that, when trusted without calibration, amplify the passive cognitive delegation that displaces the active reasoning on which durable learning depends.

Universities that deploy GenAI tools without simultaneously investing in structured AI literacy curricula and verification-centred pedagogies risk amplifying this harmful chain at the institutional scale. The non-significant moderation finding (H9: $\beta = .043$, $p = .401$) deepens the policy concern: AI literacy alone, without the deliberate cultivation of verification habits as a practised behavioural repertoire, is insufficient to interrupt the trust-induced offloading dynamic that the model identifies. Trust operates through a fast, heuristic cognitive channel that is largely impervious to deliberative metacognitive knowledge, implying that cognitive debiasing interventions are needed alongside competency-based AI literacy programming. This paradox is especially salient in the Ghanaian context, where institutional AI governance frameworks remain nascent, and student AI literacy levels are highly variable, creating a governance vacuum in which the harmful chain can operate without institutional interruption.

Policy Recommendations at Three Institutional Levels

The study's structural evidence generates concrete, evidence-based recommendations at three institutional tiers, organised from the most proximal to the most distal level of intervention.

Curriculum-Level Recommendations

At the curriculum level, assessment tasks should be redesigned to reward process-oriented AI engagement rather than merely AI-assisted product quality. Specific mechanisms include: prompt revision portfolios that require students to document and reflect on iterative AI interaction sequences, demonstrating co-creation rather than single-query acceptance; AI output audit exercises in which students systematically verify AI-generated factual claims, arguments, or analyses against authoritative sources; source triangulation requirements embedded in assignment briefs that make verification behaviour a graded performance expectation rather than an optional enhancement; and reflective AI interaction journals that require students to articulate their critical evaluation of AI outputs and document instances of AI error detection. These task designs operationalise the



co-creation ($\beta = .531$) and verification ($\beta = .540$) pathways that the HAIM identifies as the primary positive levers of educational benefit, while structurally discouraging the passive cognitive offloading (performance: 61.0%) that the model identifies as the most urgent remediation target.

Institutional-Level Recommendations

At the institutional level, Ghanaian universities should develop context-specific AI literacy frameworks that move beyond generic digital competency standards to address the specific cognitive dynamics identified by the HAIM. Such frameworks should integrate the affective, behavioural, cognitive, and ethical dimensions of AI literacy with locally relevant examples of AI limitations, common GenAI error patterns in academic contexts, and discipline-specific verification strategies. The finding that structural patterns are invariant across STEM and non-STEM groups (Table 9; all $p > .20$) strengthens the case for institution-wide rather than faculty-specific programming, reducing implementation complexity and cost. Institutions should also invest in faculty development programmes that equip academic staff to design, assess, and provide feedback on process-oriented AI engagement tasks, ensuring that the pedagogical scaffolds recommended at the curriculum level are implemented with fidelity. Additionally, institutional AI governance policies should explicitly address trust calibration: student orientation materials, AI use guidelines, and course-level briefings should communicate the risks of uncalibrated AI trust as clearly as they communicate the opportunities of AI-assisted learning.

National Policy-Level Recommendations

At the national policy level, the Ghana Tertiary Education Commission and the Ministry of Education should consider mandating minimum AI literacy competency standards for all incoming university students, analogous to digital literacy requirements now embedded in several Sub-Saharan African higher education systems. These standards should be grounded in the four dimensions of AI literacy functional understanding, prompt competency, output evaluation, and ethical awareness and assessed through standardised instruments with established validity evidence, building on the psychometric framework demonstrated in this study. National curriculum guidelines for AI-assisted learning should be developed in consultation with educational researchers, institutional leaders, and student representatives, with explicit attention to the trust-offloading paradox identified by the HAIM and the verification-centred pedagogical responses it recommends. The large predictive relevance values ($Q^2 = .389$ for learning outcomes; $Q^2 = .386$ for critical thinking) and the excellent SRMR fit (.008) confirm that the HAIM is not merely descriptively accurate but genuinely predictive, a robust empirical foundation from which longitudinal and experimental research can build to establish the causal architecture of productive human-AI interaction in Ghanaian and broader Sub-Saharan African higher education contexts.



LIMITATIONS

These findings should be interpreted with three limitations.

First, the cross-sectional design precludes causal inference; longitudinal or experimental design is needed to make directional effect and test the null moderation finding of H9 under conditions of experimentally manipulated AI literacy.

Second, all measures were based on self-report, which established the possibility of common method and social desirability biases. Though this concern is mitigated by the test of Harman (35.34%) and the procedural controls, future research should include behavioural trace data of the GenAI interaction logs to supplement self-reported patterns of interaction.

Third, although the sample is representative of the five Ghanaian institutions and balances key demographic levels, it is not representative of community-based colleges of education or technical and vocational institutions and thus should not be extrapolated to those contexts.

SUGGESTIONS FOR FUTURE RESEARCH

Four directions should be prioritised. First, longitudinal designs that would trace the development of AI literacy and the quality of interactions during an academic year would allow making causal claims that would not be possible using cross-sectional data. Second, experimental interventions with structured verification tasks represented in AI-based assignments would directly test the prescriptive implications of the HAIM on the development of critical thinking. Third, the boundary conditions would be tested in the population with different access to infrastructure in case of the extension of the HAIM to secondary education and TVET contexts in Ghana and other Sub-Saharan African countries. Fourth, experimental designs: to determine whether the null moderation is a true independence of trust and literacy effects or a cross-sectional power constraint, it would be informative to test H9 using experimental designs, in which AI literacy is manipulated by training programmes.

DECLARATIONS

Acknowledgements: Not applicable.

Funding: This research received no external funding.

Data Availability: The datasets used in this study are available from the corresponding author upon reasonable request.

Competing Interests: The authors declare no competing interests.

Author Contributions: All authors contributed equally to the development of this study.

Consent for Publication: All authors have read and approved the final manuscript.

Consent to Participate: The objectives of the study were clearly disclosed to all participants. Participation was entirely voluntary, and respondents were free to withdraw at any point. No personal identifying information was collected.

Ethics Declaration: This study did not require formal institutional review board approval. In Ghana, non-clinical social science research not involving vulnerable populations or sensitive personal data does not require formal IRB review. The study was conducted in accordance with Ghana's social science research regulations, the institutions' ethical guidelines, and the principles of the Declaration of Helsinki. This exemption determination was formally confirmed in writing by the Research and Innovation Support Directorate of the lead author's home institution, University of Skills Training and Entrepreneurial Development, Kumasi, prior to the commencement of data collection, on the basis that the study involved minimal-risk, non-clinical social science procedures with no vulnerable populations, deception, or sensitive personal data.



AI Disclosure: The authors declare that no generative AI tools, including large language models such as ChatGPT, or any other AI based writing, analysis, or content generation system, were used in conceiving the research idea, designing the study, collecting or analysing the data, or drafting, editing, or revising any part of this manuscript. Although generative AI is the subject of empirical inquiry in this study, it was not employed as a tool in producing the manuscript itself. All text, analyses, and interpretations are the original work of the named authors, who take full responsibility for the accuracy, originality, and integrity of the content presented.

REFERENCES

- 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, Article 10.
- Adewale, M. D., Azeta, A., Abayomi-Alli, A., & Sambo-Magaji, A. (2024). Empirical investigation of a multilayered framework for predicting academic performance in open and distance learning. *Electronics*, 13(14), Article 2754.
- Ahangama, N. (2026). Designing assessments in the generative AI era: A tailored assessment framework for ICT tertiary education. *International Journal of Educational Technology in Higher Education*, 23(1). <https://doi.org/10.1186/s41239-026-00582-0>
- Alshamy, A., Al-Harthi, A. S. A., & Abdullah, S. (2025). Perceptions of Generative AI Tools in Higher Education: Insights from Students and Academics at Sultan Qaboos University. *Education Sciences*, 15(4), 501. <https://doi.org/10.3390/educsci15040501>
- Asghar, M. Z., Duah, K. A., Iqbal, J., & Järvenoja, H. (2025). Evidence from West Africa on the interplay of affective, behavioural, cognitive, and ethical dimensions of AI literacy in Ghanaian and Nigerian Universities. *Discover Computing*, 28(1). <https://doi.org/10.1007/s10791-025-09691-2>
- Ayanwale, M. A., Adelana, O. P., Bamiro, N., Olatunbosun, S., Idowu, K. O., & Adewale, K. A. (2025). Large language models and GenAI in education: Insights from Nigerian in-service teachers through a hybrid ANN-PLS-SEM approach. *F1000Research*, 14, 101.
- Chatterjee, S., & Bhattacharjee, K. (2020). Adoption of artificial intelligence in higher education: A quantitative analysis using structural equation modelling. *Education and Information Technologies*, 25(4), 3443–3463.
- Chen, K., Tallant, A. C., & Selig, I. (2024). Exploring generative AI literacy in higher education: student adoption, interaction, evaluation and ethical perceptions. *Information and Learning Sciences*, 126(1/2), 132–148. <https://doi.org/10.1108/ils-10-2023-0160>
- Chiu, T. K. (2023). Future research recommendations for transforming higher education with generative AI. *Computers and Education Artificial Intelligence*, 6, 100197. <https://doi.org/10.1016/j.caeai.2023.100197>
- Chiu, T. K. F., Ahmad, Z., Ismailov, M., & Sanusi, I. T. (2024). What are artificial intelligence literacy and competency? A comprehensive framework to support them. *Computers and Education Open*, 5, 100154.
- Clark, A., & Chalmers, D. (1998). The extended mind. *Analysis*, 58(1), 7–19. <https://doi.org/10.1093/analys/58.1.7>
- Crompton, H., Burke, D., & Nickel, C. (2026). Designing faculty standards for technology integration in higher education institutions: a design-based research study. *International Journal of Educational Technology in Higher Education*, 23(1). <https://doi.org/10.1186/s41239-026-00584-y>
- Dong, H., Chen, X., & Shen, J. (2026). Unpacking the impact mechanism of generative AI-based learning analytics technology on student learning engagement. *International Journal of Educational Technology in Higher Education*, 23(1). <https://doi.org/10.1186/s41239-026-00591-z>
- Elycheikh, A., Svetlana, M., & Magda, P. (2024). Critical Integration of Generative AI in Higher Education: Cognitive, Pedagogical and Ethical Perspectives. *Global Journal of Human-Social Science*, 1–12. <https://doi.org/10.34257/ljrjssvol25is13pg1>



- Essel, H. B., Vlachopoulos, D., Essuman, A. B., & Amankwa, J. O. (2023). ChatGPT effects on cognitive skills of undergraduate students: Receiving instant responses from AI-based conversational large language models (LLMs). *Computers and Education Artificial Intelligence*, 6, 100198. <https://doi.org/10.1016/j.caeai.2023.100198>
- 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>
- Farrokhnia, M., Latifi, S., Papadopoulos, P. M., Hogenkamp, L., Gijlers, H., Khosravi, H., & Noroozi, O. (2026). Generative AI offers more, but students revise less: comparing the effects of teacher and AI feedback on student essay revisions. *International Journal of Educational Technology in Higher Education*, 23(1). <https://doi.org/10.1186/s41239-026-00579-9>
- Fornell, C., & Larcker, D. F. (1981). Evaluating structural equation models with unobservable variables and measurement error. *Journal of Marketing Research*, 18(1), 39–50.
- Francis, N. J., Jones, S., & Smith, D. P. (2025). Generative AI in Higher Education: Balancing innovation and integrity. *British Journal of Biomedical Science*, 81, 14048. <https://doi.org/10.3389/bjbs.2024.14048>
- Gao, Z., Cheah, J., Lim, X., & Luo, X. (2024). Enhancing academic performance of business students using generative AI: An interactive-constructive-active-passive (ICAP) self-determination perspective. *The International Journal of Management Education*, 22(1), 100896.
- Gerlich, M. (2025). AI tools in society: Impacts on cognitive offloading and the future of critical thinking. *Societies*, 15(1), Article 12.
- Goto, J., & Ramnarain, U. (2026). Does culture matter in AI adoption? A predictive analysis of cultural dimensions, social influence, and personal innovativeness in the modified UTAUT model. *International Journal of Educational Technology in Higher Education*, 23(1). <https://doi.org/10.1186/s41239-026-00588-8>
- Granć, A. (2025). Emerging Drivers of Adoption of Generative AI Technology in Education: A Review. *Applied Sciences*, 15(13), 6968. <https://doi.org/10.3390/app15136968>
- Grinschgl, S., & Neubauer, A. (2022). Supporting cognition with modern technology: Distributed cognition today and in an AI-enhanced future. *Frontiers in Artificial Intelligence*, 5, 918321.
- Hair, J. F., Risher, J. J., Sarstedt, M., & Ringle, C. M. (2021). When to use and how to report the results of PLS-SEM. *European Business Review*, 31(1), 2–24.
- Haroud, S., & Saqri, N. (2025). Generative AI in Higher Education: Teachers' and students' perspectives on support, replacement, and digital literacy. *Education Sciences*, 15(4), 396. <https://doi.org/10.3390/educsci15040396>
- Hazaimah, M., & Al-Ansi, A. (2024). Model of AI acceptance in higher education: Arguing teaching staff and students' perspectives. *The International Journal of Information and Learning Technology*, 41(5), 512–530.
- Henseler, J., Ringle, C. M., & Sarstedt, M. (2015). A new criterion for assessing discriminant validity in variance-based structural equation modelling. *Journal of the Academy of Marketing Science*, 43(1), 115–135.
- Hong, H., Vate-U-Lan, P., & Viriyavejakul, C. (2025). Cognitive Offload Instruction with Generative AI: A Quasi-Experimental Study on Critical Thinking Gains in English Writing. *Forum for Linguistic Studies*, 7(7). <https://doi.org/10.30564/fls.v7i7.10072>
- Hu, X., Luo, L., & Fleming, S. M. (2019). A role for metamemory in cognitive offloading. *Cognition*, 193, 104012.
- Iqbal, J., Hashmi, Z. F., Asghar, M. Z., & Abid, M. N. (2025). Generative AI tool use enhances academic achievement in sustainable education through shared metacognition and cognitive offloading among preservice teachers. *Scientific Reports*, 15(1), 16610. <https://doi.org/10.1038/s41598-025-01676-x>
- Kim, J., Lee, S., Detrick, R., Wang, J., & Li, N. (2025). Students' Generative AI interaction patterns and their impact on academic writing. *Journal of Computing in Higher Education*, 38(1), 504–525. <https://doi.org/10.1007/s12528-025-09444-6>



- Liu, Z., Zuo, H., & Lu, Y. (2025). The impact of ChatGPT on students' academic achievement: A meta-analysis. *Journal of Computer Assisted Learning*, 41(4), 987–1006.
- Mahama, I., & Amadu, P. (2025). You are the Driver, and AI is the Mate: Exploring Human-Led Creative and Critical Thinking in AI-Augmented Learning Environments. *F1000Research*, 14, 974. <https://doi.org/10.12688/f1000research.167988.1>
- Mante, D. A., Opoku, O. G., Chineta, O. M., Adamu, A., Parker, D., Xuefu, X., Hui, X., Setsoafia, N. K., Kyeremeh, P., & Opoku, D. (2025). Examining the moderating role of generative AI in advancing teachers' professional competencies through self-directed learning: evidence from Ghana. *Interactive Learning Environments*, 1–17. <https://doi.org/10.1080/10494820.2025.2556811>
- Naamati-Schneider, L., & Alt, D. (2024). Beyond digital literacy: The era of AI-powered assistants and evolving user skills. *Education and Information Technologies*, 29(5), 6123–6147.
- Ng, D. T. K., Leung, J. K. L., Su, J. K. L., Ng, R., & Chu, S. K. W. (2023). Teachers' AI digital competencies and twenty-first-century skills in the post-pandemic world. *Educational Technology Research and Development*, 71(1), 45–66.
- Ng, D. T. K., Wu, W., Leung, J. K. L., Chiu, T. K. F., & Chu, S. K. W. (2024). Design and validation of the AI literacy questionnaire: The affective, behavioural, cognitive and ethical approach. *British Journal of Educational Technology*, 55(4), 1082–1104.
- Podsakoff, P. M., MacKenzie, S. B., Lee, J.-Y., & Podsakoff, N. P. (2003). Common method biases in behavioural research: A critical review of the literature and recommended remedies. *Journal of Applied Psychology*, 88(5), 879–903.
- Qian, Y. (2025). Pedagogical Applications of Generative AI in Higher Education: A Systematic Review of the Field. *TechTrends*, 69(5), 1105–1120. <https://doi.org/10.1007/s11528-025-01100-1>
- Risko, E. F., & Gilbert, S. J. (2016). Cognitive offloading. *Trends in Cognitive Sciences*, 20(9), 676–688.
- Ruiz-Rojas, L. I., Salvador-Ullauri, L., & Acosta-Vargas, P. (2024). Collaborative working and critical Thinking: adoption of generative artificial intelligence tools in higher education. *Sustainability*, 16(13), 5367. <https://doi.org/10.3390/su16135367>
- Saihi, A., & Ahmed, V. (2026). Uncovering adoption personas for generative AI in higher education: a clustering-based segmentation approach. *International Journal of Educational Technology in Higher Education*, 23(1). <https://doi.org/10.1186/s41239-026-00583-z>
- Salifu, I., Arthur, F., Arkorful, V., Nortey, S. A., & Osei-Yaw, R. S. (2024). Economics students' behavioural intention and usage of ChatGPT in higher education: A hybrid structural equation modelling-artificial neural network approach. *Cogent Social Sciences*, 10(1), 2301234.
- Sergeeva, O. V., Zheltukhina, M. R., Shoustikova, T., Tukhvatullina, L. R., Dobrokhoto, D. A., & Kondrashev, S. V. (2025). Understanding higher education students' adoption of generative AI technologies: An empirical investigation using UTAUT2. *Contemporary Educational Technology*, 17(2), ep571. <https://doi.org/10.30935/cedtech/16039>
- Shahzad, M. F., Xu, S., & Zahid, H. (2024). Exploring the impact of generative AI-based technologies on learning performance through self-efficacy, fairness & ethics, creativity, and trust in higher education. *Education and Information Technologies*, 30(3), 3691–3716. <https://doi.org/10.1007/s10639-024-12949-9>
- Skulmowski, A. (2023). The cognitive architecture of digital externalisation. *Educational Psychology Review*, 35(3), 1647–1673.
- Strzelecki, A., & ElArabawy, S. (2024). Investigation of the moderation effect of gender and study level on the acceptance and use of generative AI by higher education students. *British Journal of Educational Technology*, 55(3), 1209–1230. <https://doi.org/10.1111/bjet.13425>



- Tzirides, A. O., Zapata, G., Kastania, N. P., Saini, A. K., Castro, V., Ismael, S. A., You, Y., Santos, T. A. D., Searsmith, D., O'Brien, C., Cope, B., & Kalantzis, M. (2024). Combining human and artificial intelligence for enhanced AI literacy in higher education. *Computers and Education Open*, 6, 100184. <https://doi.org/10.1016/j.caeo.2024.100184>
- Wang, S., & Zhang, H. (2026a). Pedagogical partnerships with generative AI in higher education: how dual cognitive pathways paradoxically enable transformative learning. *International Journal of Educational Technology in Higher Education*, 23(1). <https://doi.org/10.1186/s41239-026-00585-x>
- Weis, L., Bele, J. L., & Erčulj, V. (2026). Acceptance and Use of Generative Artificial Intelligence in Higher Education: A UTAUT-Based model integrating trust and privacy. *Education Sciences*, 16(2), 173. <https://doi.org/10.3390/educsci16020173>
- Wiredu, J. K., Abuba, N. S., & Zakaria, H. (2024). Impact of generative AI on academic integrity and learning Outcomes: a case study in the Upper East region. *Asian Journal of Research in Computer Science*, 17(8), 70–88. <https://doi.org/10.9734/ajrcos/2024/v17i7491>
- Wu, D., & Zhang, J. (2025). Generative artificial intelligence in secondary education: Applications and effects on students' innovation skills and digital literacy. *PLoS ONE*, 20(5), e0323349. <https://doi.org/10.1371/journal.pone.0323349>
- Yakubu, M. N., David, N., & Abubakar, N. H. (2025). Students' behavioural intention to use content generative AI for learning and research: A UTAUT theoretical perspective. *Education and Information Technologies*, 30(13), 17969–17994. <https://doi.org/10.1007/s10639-025-13441-8>
- Yilmaz, M., Temur, H. B., Emmungil, L., Çelik, E., Gauthier, A., & Cukurova, M. (2026). Supporting self-regulated learning through generative AI feedback in online higher education: the importance of student perceptions of the source of feedback. *International Journal of Educational Technology in Higher Education*, 23, 16. <https://doi.org/10.1186/s41239-026-00592-y>
- Zaim, M., Arsyad, S., Waluyo, B., Ardi, H., Hafizh, M. A., Zakiyah, M., Syafitri, W., Nusi, A., & Hardiah, M. (2025). Generative AI as a cognitive Co-Pilot in English language learning in higher education. *Education Sciences*, 15(6), 686. <https://doi.org/10.3390/educsci15060686>
- Zhao, Z., An, Q., & Liu, J. (2025). Exploring AI tool adoption in higher education: evidence from a PLS-SEM model integrating multimodal literacy, self-efficacy, and university support. *Frontiers in Psychology*, 16, 1619391. <https://doi.org/10.3389/fpsyg.2025.1619391>
- Zimmerman, B. J. (2002). Becoming a self-regulated learner: An overview. *Theory Into Practice*, 41(2), 64–70.