



Critical Evaluation of AI-Generated Academic Information in Higher Education: Empirical Evidence on Trust, AI-Related Digital Literacy and Verification Behaviour

Dr. Kuldeep Kaur¹, Prof. Pramendra Singh Pundir²

¹Academic Researcher and Educator, Ph.D., Department of Education, University of Allahabad, India

¹Email: kuldeep5bioinfo@gmail.com

²Professor and Head, Department of Statistics, University of Allahabad, India

ABSTRACT

Generative artificial intelligence (GenAI) is increasingly used by university students for academic information seeking, explanation, drafting and problem solving. Its usefulness, however, depends not only on ease of use or frequency of interaction but also on users' capacity to evaluate and verify AI-generated information. This study presents an independent secondary-data analysis of a publicly available, anonymised dataset containing responses from 104 university students to 20 five-point Likert-scale items covering perceived ease of use (PEOU), trust in AI, verification behaviour and digital literacy. The analysis examined internal consistency, factorability diagnostics, descriptive statistics, Pearson correlations and multiple linear regression using heteroscedasticity-consistent HC3 standard errors. The four constructs demonstrated acceptable to high internal consistency (Cronbach's $\alpha = .782-.898$). The KMO measure was .836 and Bartlett's test of sphericity was statistically significant, $\chi^2(190) = 1087.35, p < .001$, indicating that the item correlation matrix was suitable for factor-analytic examination. Verification behaviour was positively correlated with AI-related digital literacy ($r = .445, p < .001$), whereas its bivariate associations with PEOU and trust were not statistically significant. In the multivariable model, AI-related digital literacy showed the largest standardized coefficient ($\beta = .524, p < .001$). Trust showed a small negative adjusted association with verification behaviour ($\beta = -.214, p = .048$), whereas PEOU and usage frequency were not statistically significant. The model explained 25.3% of the variance in verification behaviour ($R^2 = .253$; adjusted $R^2 = .223$). The findings suggest that effective academic engagement with GenAI should not be equated with frequent use or perceived ease of use. Instead, evaluative digital-literacy capabilities appear particularly relevant to verification behaviour. Building on these findings, the study proposes the HEAVEN framework (Human-centred Evaluation of AI-generated academic information through Verification, Evidence and Nuanced judgement) as a conceptual framework integrating AI literacy, information verification, evidence appraisal and human judgement in higher education.

Keywords: Generative artificial intelligence; higher education; AI-related digital literacy; verification behaviour; trust in AI; HEAVEN framework

1. Introduction

Generative artificial intelligence (GenAI) has rapidly become part of the academic information environment. Systems capable of producing natural-language responses can assist students with explanation, summarisation, brainstorming, drafting, coding, translation and information seeking. Their accessibility and conversational interface make them particularly attractive for educational use. At the same time, the fluency of AI-generated responses can create an important epistemic problem: an answer may appear coherent and authoritative even when its claims are incomplete, unsupported, outdated or incorrect. This issue changes the nature of academic information literacy. Traditional

information-literacy practices emphasise identifying information needs, locating sources, evaluating credibility and using information responsibly. In an AI-mediated environment, these practices must be extended to situations in which the first information encountered may itself be synthetically generated. Users therefore need to distinguish between the ability to obtain an answer from an AI system and the ability to determine whether that answer should be accepted, qualified or rejected.

Research on artificial intelligence in higher education has highlighted both the educational opportunities and the risks associated with generative systems. Kasneci et al. (2023), for example, discuss the potential of large language models for learning while also identifying concerns relating to accuracy, bias and responsible use. More broadly, research on artificial intelligence in education has emphasised the importance of human judgement, pedagogical context and institutional conditions (Zawacki-Richter et al., 2019; Lo, 2023; Zhai et al., 2024). UNESCO guidance similarly places human agency, critical engagement and responsible use at the centre of generative-AI adoption in education (Miao & Holmes, 2023). A central issue is therefore not simply whether students use AI but how they respond to AI-generated information. A student may use an AI system frequently and find it easy to operate while making little effort to verify the resulting claims. Conversely, a student may use AI as a preliminary information tool while systematically checking important claims against external evidence. These are substantively different forms of AI engagement. Trust is particularly relevant to this distinction. Some degree of trust is necessary for productive interaction with AI systems; however, reliance that is insufficiently calibrated to system limitations may reduce the motivation to question or verify generated outputs. Responsible AI use therefore requires not the elimination of trust but its calibration to the quality, context and risk of the information being used. Digital literacy may be especially important in this process. In the present study, AI-related digital literacy refers to the capacity to engage with AI-mediated information in an informed and evaluative manner. This includes practical interaction with AI systems as well as the ability to recognise limitations, inspect outputs, locate supporting evidence and make informed decisions about whether information can appropriately be used. The distinction between AI-use fluency and verification competence is consequently central to understanding responsible academic use of GenAI.

The present study addresses this issue through an independent secondary-data analysis of a publicly available dataset containing responses from 104 university students who reported using AI chatbots for academic activities. The source dataset was collected and deposited by Farras, Lusiawati and Simatupang (2026). Rather than treating the present study as a new data-collection exercise, the analysis focuses specifically on verification behaviour as the central outcome and examines its associations with PEOU, trust in AI, AI-related digital literacy and usage frequency. The study makes three related contributions. First, it provides an independent secondary-data statistical analysis using a focused research question centred on verification behaviour. Second, it examines whether ease of use, trust, AI-related digital literacy and frequency of use show similar or different statistical relationships with verification behaviour within a multivariable framework. Third, it develops the HEAVEN framework, a theoretically derived and empirically informed framework for human-centred evaluation of AI-generated academic information. The study does not claim ownership of the original observations or the original data-collection process. The original dataset creators are explicitly acknowledged while the present article's contribution lies in its analytical framing, statistical interpretation, theoretical synthesis and proposed framework.

2. Conceptual Background

2.1 Generative AI and Academic Information Use

Generative AI systems have introduced a new form of information interaction in higher education. Instead of navigating a set of documents and evaluating individual sources before forming an answer, users can request a synthesised response directly from an AI system. This can reduce the time and effort required to obtain information but it can also obscure the provenance of individual claims. The problem is particularly significant in academic contexts. Students may use AI-generated content in assignments, presentations, literature reviews, discussions or independent learning. When generated information is inaccurate, fabricated or inadequately supported, users who accept it without verification can inadvertently transfer these weaknesses into academic work. Consequently, the educational objective should not be limited to teaching students how to operate AI tools. It should also involve developing the capacity to question, verify, contextualise and appropriately communicate AI-generated information.

2.2 Information Literacy, Critical Thinking and Verification

Information literacy involves the ability to recognise information needs, locate relevant information, evaluate sources and use information responsibly. The Framework for Information Literacy for Higher Education emphasises practices such as authority evaluation, information creation as a process and strategic exploration (Association of College and Research Libraries [ACRL], 2016). AI-generated information introduces an additional evaluative layer. A generated response may contain several claims but those claims do not automatically possess the evidentiary status of a scholarly source. Verification therefore involves identifying important claims and checking them against appropriate external evidence. Critical thinking is closely related to this process. Facione (1990) conceptualises critical thinking in terms of skills such as interpretation, analysis, evaluation, inference, explanation and self-regulation. In an AI-mediated academic environment, these capabilities can be applied to AI-generated outputs by asking whether a claim is adequately supported, whether alternative explanations exist and whether the available evidence justifies the conclusion. The present study therefore treats verification behaviour as a behavioural component of responsible AI engagement rather than as a synonym for general AI use.

2.3 Trust and Reliance on AI

Trust can facilitate interaction with AI systems because users need some degree of confidence before relying on system outputs. However, trust that is not calibrated to system capability may encourage over-reliance. This distinction is important because AI systems can generate persuasive language without guaranteeing factual accuracy. Responsible use consequently requires users to recognise that confidence in the interaction interface is not equivalent to confidence in the truth of an individual claim.

The present analysis does not assume that trust necessarily causes either greater or lower verification. Instead, it examines whether reported trust is statistically associated with verification behaviour after accounting for PEOU, AI-related digital literacy and usage frequency.

2.4 AI-Related Digital Literacy

For the purposes of this study, AI-related digital literacy refers to users' capacity to engage with AI-mediated information in an informed and evaluative manner. It extends conventional digital literacy by incorporating awareness of AI-generated content, its limitations and the need for appropriate verification. This distinction is important because technical familiarity with an AI interface does not necessarily imply evaluative competence. A user may be highly comfortable with prompting and interacting with an AI system while having limited capacity to assess the reliability of its output.

The empirical analysis therefore separates PEOU from AI-related digital literacy. PEOU captures perceived usability, whereas AI-related digital literacy is interpreted as a broader evaluative capability.

3. Research Gap and Objectives

Existing discussions of generative AI in higher education frequently address adoption, usefulness, student attitudes, academic integrity and technological acceptance. Comparatively less attention has been given to the behavioural question of which user characteristics are associated with students' verification of AI-generated academic information.

Three issues are particularly relevant.

First, frequent AI use is not necessarily equivalent to responsible AI use. Usage frequency may indicate exposure but does not establish whether users evaluate generated outputs.

Second, perceived ease of use may facilitate interaction with an AI system but does not necessarily indicate information-evaluation competence.

Third, trust may have a complex relationship with verification. Trust may facilitate productive interaction while uncritical reliance may reduce the perceived need for independent checking.

Accordingly, this study has four objectives:

- (i) To assess the internal consistency and factorability diagnostics of the four constructs in the secondary dataset.
- (ii) To examine the associations among PEOU, trust in AI, AI-related digital literacy and verification behaviour.
- (iii) To estimate the independent associations of PEOU, trust, AI-related digital literacy and usage frequency with verification behaviour.

- (iv) To develop a human-centred conceptual framework for evaluating AI-generated academic information based on the empirical findings.

The corresponding research questions are:

RQ1: What are the descriptive and internal-consistency characteristics of PEOU, trust in AI, verification behaviour and AI-related digital literacy?

RQ2: How are PEOU, trust in AI, AI-related digital literacy and verification behaviour related to one another?

RQ3: Which of PEOU, trust, AI-related digital literacy and usage frequency show statistically significant independent associations with verification behaviour?

RQ4: How can the empirical findings inform a human-centred framework for evaluating AI-generated academic information?

4. Method

4.1 Research Design

The study uses a quantitative secondary-data research design. The observations were not collected by the present authors. Instead, the analysis uses a publicly available, anonymised dataset deposited in Zenodo by Farras, Lusiawati and Simatupang (2026). The source record reports survey data from 104 university students who used AI chatbots for academic activities. The dataset contains 20 five-point Likert-scale items covering four constructs: PEOU, Trust in AI, Verification Behaviour and Digital Literacy.

The present study constitutes an independent analytical contribution. The original dataset is acknowledged as the source of the observations, whereas the present authors formulated the present research focus, specified the statistical model, conducted the secondary statistical analysis, interpreted the findings in relation to AI-related digital literacy and verification and developed the HEAVEN framework. As the source dataset was designed to support multiple statistical analyses, the present article does not claim that the dataset itself is new. Instead, its contribution is positioned in terms of the present study's focused analytical question, multivariable interpretation and conceptual framework.

4.2 Data Source

The dataset used in the study is: Farras, A. J., Lusiawati, D. A. & Simatupang, S. S. S. (2026). Survey Dataset for Exploring Students' Trust and Verification Behaviour in AI Chatbot Use [Data set]. Zenodo. doi:10.5281/zenodo.21135117. The dataset is publicly accessible and anonymised. According to the repository description, personally identifiable information was removed. The original sampling procedure was purposive and the dataset contains 104 student-level observations. No new participants were recruited for the present study.

4.3 Measures

The source instrument consists of 20 five-point Likert-scale items distributed across four constructs.

Construct	Number of items	Analytical role
Perceived Ease of Use (PEOU)	5	Perceived usability of AI systems
Trust in AI	5	Confidence/reliance in AI-generated outputs
Verification Behaviour	5	Behavioural tendency to check AI-generated information
Digital Literacy	5	AI-related evaluative and digital capability

For consistency with the present theoretical framing, the source construct labelled Digital Literacy is referred to throughout this article as AI-related digital literacy. This relabelling is interpretive and does not imply that a new instrument was developed. Composite construct scores were represented by the mean of the respective five items.

4.4 Descriptive Characteristics of the Sample

The dataset contains 104 university students. Semester distribution was as follows: Semester 2, 14 students (13.5%); Semester 4, 73 (70.2%); Semester 6, 10 (9.6%); and Semester 8, 7 (6.7%). Reported AI-system use was distributed across ChatGPT, 64 students (61.5%); Claude, 24 (23.1%); Gemini, 14 (13.5%); and DeepSeek, 2 (1.9%).

Usage-frequency values were recorded in the source data as categories 3, 4 and 5, with frequencies of 9 (8.7%), 43 (41.3%) and 52 (50.0%), respectively. These numerical categories are retained as coded in the source data. Their substantive labels should be interpreted only with reference to the original coding scheme.

4.5 Statistical Analysis

The analysis proceeded in five stages. First, descriptive statistics were calculated for the four composite constructs, including means, standard deviations and observed ranges. Second, internal consistency was assessed using Cronbach's alpha. These coefficients are interpreted as evidence of internal consistency rather than as evidence of complete psychometric validation.

Third, KMO and Bartlett's test of sphericity were examined as factorability diagnostics. Since, the purpose of this study was construct-level association and prediction rather than development or validation of a new measurement instrument, exploratory factor analysis was not treated as a primary analytical objective. The factorability results are therefore reported as supplementary diagnostics. Fourth, Pearson product-moment correlations were calculated to examine bivariate associations among the four constructs. Finally, multiple linear regression was used with verification behaviour as the dependent variable and PEOU, trust, AI-related digital literacy and usage frequency as predictors. Heteroscedasticity-consistent HC3 standard errors were used for coefficient inference. The cross-sectional design means that the resulting coefficients are interpreted as statistical associations rather than causal effects. The statistical analysis was conducted using Python-based statistical procedures.

5. Results

5.1 Descriptive Statistics and Internal Consistency

The four constructs showed generally favourable mean scores, although trust was lower than the other three constructs.

Construct	M	SD	Minimum	Maximum	Cronbach's α
PEOU	4.065	0.613	1.80	5.00	.795
Trust in AI	3.406	0.812	1.00	5.00	.898
Verification Behaviour	4.271	0.610	2.20	5.00	.849
AI-related Digital Literacy	3.948	0.594	2.60	5.00	.782

The internal-consistency coefficients ranged from .782 to .898. These values provide evidence that the items within each construct showed acceptable to high internal consistency in this sample. They should not, however, be interpreted as evidence that the scales are fully validated across populations.

Verification behaviour had the highest mean score ($M = 4.271$, $SD = 0.610$), followed by PEOU ($M = 4.065$, $SD = 0.613$) and AI-related digital literacy ($M = 3.948$, $SD = 0.594$). Trust in AI had a comparatively lower mean ($M = 3.406$, $SD = 0.812$) and greater dispersion.

5.2 Factorability Diagnostics

The KMO measure of sampling adequacy was .836, indicating that the correlation structure was suitable for factor-analytic examination. Bartlett's test of sphericity was statistically significant, $\chi^2(190) = 1087.35$, $p < .001$.

These diagnostics indicate that the item correlation matrix contains systematic relationships rather than behaving like an identity matrix. Because the present research does not seek to establish a new factor structure, these statistics are treated as supplementary evidence rather than as psychometric validation.

5.3 Correlations Among Study Variables

Pearson correlations are presented below:

Variables	r	p
PEOU – Trust	.641	< .001
PEOU – Verification Behaviour	.008	.934
PEOU – AI-related Digital Literacy	.374	< .001
Trust – Verification Behaviour	-.086	.384
Trust – AI-related Digital Literacy	.300	.002
Verification Behaviour – AI-related Digital Literacy	.445	< .001

PEOU was positively associated with trust in AI ($r = .641$, $p < .001$) and with AI-related digital literacy ($r = .374$, $p < .001$). However, PEOU had essentially no bivariate association with verification behaviour ($r = .008$, $p = .934$).

Trust was positively associated with AI-related digital literacy ($r = .300$, $p = .002$) but was not significantly associated with verification behaviour at the bivariate level ($r = -.086$, $p = .384$).

The largest correlation involving verification behaviour was its positive association with AI-related digital literacy ($r = .445$, $p < .001$). This finding provides the basis for examining AI-related digital literacy more closely in the multivariable model.

5.4 Multiple Linear Regression

Verification behaviour was modelled as the dependent variable. PEOU, trust in AI, AI-related digital literacy and usage frequency were entered as predictors.

The model explained 25.3% of the variance in verification behaviour ($R^2 = .253$; adjusted $R^2 = .223$). Coefficient inference used HC3 heteroscedasticity-consistent standard errors.

Predictor	B	HC3 SE	Standardized β	p	95% CI
Intercept	2.838	.492	—	< .001	[1.873, 3.803]
PEOU	-.076	.122	-.077	.531	[-.315, .163]
Trust in AI	-.161	.081	-.214	.048	[-.320, -.002]
AI-related Digital Literacy	.538	.132	.524	< .001	[.279, .796]
Usage frequency	.038	.107	.040	.724	[-.173, .249]

AI-related digital literacy had the largest standardized coefficient in the model ($\beta = .524$, $p < .001$). Holding the other predictors constant, higher AI-related digital-literacy scores were associated with higher verification-behaviour scores.

Trust in AI showed a small negative adjusted association with verification behaviour ($\beta = -.214$, $p = .048$). Notably, the corresponding bivariate correlation was not statistically significant. The difference between the bivariate and adjusted estimates indicates that the multivariable coefficient should not be interpreted as a simple bivariate relationship.

PEOU was not statistically significant ($\beta = -.077$, $p = .531$) and usage frequency was also not statistically significant ($\beta = .040$, $p = .724$).

The regression findings therefore distinguish between using AI comfortably or frequently and engaging in verification behaviour. In this sample, AI-related digital literacy showed the largest standardized association with verification behaviour among the predictors examined.

6. Discussion

6.1 AI-Related Digital Literacy and Verification

The central finding is the positive relationship between AI-related digital literacy and verification behaviour. The bivariate association was moderate in magnitude ($r = .445$, $p < .001$) and AI-related digital literacy retained the largest standardized coefficient in the multivariable regression ($\beta = .524$, $p < .001$). This result is consistent with the broader

conceptual distinction between AI-use fluency and evaluative competence. A student may know how to access, prompt and interact with an AI system without necessarily possessing the skills required to determine whether its output is reliable.

The finding therefore supports an educational interpretation in which AI literacy should include more than technical interaction. Students need opportunities to identify claims, inspect sources, compare information, recognise uncertainty and determine whether a generated response is appropriate for academic use. Significantly, the result is an association rather than evidence of causality. The cross-sectional design cannot establish whether increasing AI-related digital literacy would itself produce greater verification behaviour.

6.2 Perceived Ease of Use Is Not Equivalent to Verification Competence

PEOU showed a strong positive association with trust ($r = .641, p < .001$), suggesting that students who perceived AI systems as easier to use also tended to report greater trust. However, PEOU had virtually no association with verification behaviour ($r = .008, p = .934$) and was not statistically significant in the multivariable model. This distinction has an important educational implication. Making AI systems easier to access may improve user experience but usability alone does not guarantee responsible information evaluation.

The findings therefore suggest that successful AI adoption should not automatically be treated as equivalent to successful AI literacy. A student who can use an AI system efficiently may still require explicit instruction in source evaluation and verification.

6.3 Trust and Verification

The relationship between trust and verification requires careful interpretation. At the bivariate level, trust was not significantly related to verification behaviour ($r = -.086, p = .384$). After adjustment for PEOU, AI-related digital literacy and usage frequency, however, trust showed a small negative coefficient ($\beta = -.214, p = .048$). This pattern is consistent with the possibility that, after accounting for other measured characteristics, higher trust may be associated with lower reported verification. However, the analysis does not establish that trust causes reduced verification.

The coefficient is relatively small and statistically close to the conventional .05 threshold. It should therefore be interpreted as a small negative adjusted association, rather than evidence that trust inherently reduces verification. Future research should directly examine whether calibrated trust, perceived AI reliability, reliance and verification interact in more complex ways.

6.4 Usage Frequency and Responsible AI Engagement

Usage frequency was not statistically significant in the multivariable model ($\beta = .040, p = .724$). This indicates that frequency of interaction with AI systems was not independently associated with verification behaviour after accounting for the other variables in the model.

This finding reinforces an important distinction: exposure is not competence. Frequent use may increase familiarity with AI systems but familiarity does not necessarily result in better information evaluation. From an educational perspective, this means that simply encouraging or monitoring AI use is insufficient. Institutions may obtain more meaningful outcomes by embedding verification tasks directly into AI-supported academic activities.

6.5 From AI Adoption to Epistemic Responsibility

Taken together, the findings indicate that responsible AI engagement involves an epistemic dimension. Students must not only know how to use AI systems but also know how to assess the information they generate. This perspective is consistent with responsible-AI approaches that place human judgement and contextual evaluation alongside technical capability. Recent work by Kaur and Pundir (2026), for example, emphasises professional judgement and institutional conditions in responsible AI readiness in higher education. The present study extends this perspective to the student level by focusing specifically on verification behaviour.

7. The Proposed HEAVEN Framework

Based on the empirical findings and the conceptual integration of AI literacy, information literacy, evidence evaluation and critical judgement, this study proposes the following framework:

HEAVEN (Human-centred Evaluation of AI-generated academic information through Verification, Evidence and Nuanced judgement)

The framework is intended as a proposed, theoretically derived and empirically informed framework, not as a validated measurement model. Its purpose is to provide a practical structure through which students can evaluate AI-generated academic information.

7.1 H — Human-centred Evaluation

The human user remains responsible for the final academic judgement. AI-generated text should be treated as an input to reasoning rather than an autonomous authority.

The key question is:

What should I, as the academic user, decide after examining this AI-generated information?

7.2 E — Evaluation of Evidence

Users should identify the evidence supporting important AI-generated claims. Where possible, claims should be traced to authoritative, relevant and current sources.

The key question is:

What evidence supports this claim?

7.3 A — Assessment of AI Output

AI-generated outputs should be assessed for accuracy, relevance, completeness, contextual appropriateness and potential risk.

The key question is:

How reliable and appropriate is this output for the specific academic purpose?

7.4 V — Verification

Important claims should be independently checked against appropriate external sources. Verification should be proportionate to the consequences of error.

The key question is:

Can this claim be independently confirmed?

7.5 E — Evidence-based Reasoning

The user should distinguish between what the AI system states, what the available evidence demonstrates and what can reasonably be inferred.

The key question is:

Does the evidence justify the conclusion being presented?

7.6 N — Nuanced Judgement

Academic decisions should recognise uncertainty, limitations, competing interpretations and contextual differences.

The key question is:

What is justified, what remains uncertain and how should that uncertainty be communicated?

7.7 Conceptual Flow

The framework can be represented as:

AI-generated information → Human-centred evaluation → Evidence assessment → Output assessment → Verification → Evidence-based reasoning → Nuanced academic judgement

The framework does not imply that every AI-generated statement requires identical levels of scrutiny. Verification intensity should depend on the importance and potential consequences of the claim. For example, a low-risk brainstorming suggestion may require limited checking, whereas a medical, legal, statistical, scientific or assessment-related claim may require substantially stronger evidence verification. The HEAVEN framework should therefore be understood as a conceptual decision framework, rather than as a claim that each dimension has already been independently measured and validated.

8. Theoretical Contributions

8.1 Distinguishing AI-Use Fluency from Verification Competence

The findings contribute to the conceptual distinction between technological interaction and epistemic competence. PEOU and usage frequency represent dimensions of interaction with AI, whereas verification behaviour represents an

evaluative response to AI-generated information. The absence of a significant relationship between PEOU and verification, together with the positive association between AI-related digital literacy and verification, supports treating these as conceptually distinct dimensions.

8.2 Integrating AI Literacy and Information Literacy

The study brings AI-related digital literacy into closer conceptual connection with established information-literacy practices. AI literacy should not be limited to prompt formulation or tool familiarity. It should also include claim evaluation, source checking, evidence appraisal, recognition of uncertainty and responsible judgement.

8.3 Calibrated Trust

The adjusted negative association between trust and verification suggests that trust deserves more nuanced treatment in future research. The relevant educational objective may not be simply increasing or decreasing trust but developing calibrated trust; trust that is responsive to context, evidence, system limitations and the consequences of error.

8.4 HEAVEN as a Human-Centred Evaluation Model

The proposed HEAVEN framework provides a conceptual bridge between AI literacy, information literacy, verification behaviour, evidence evaluation and responsible academic judgement. Its contribution is therefore primarily conceptual at this stage. It should not be described as a validated framework until its dimensions and measurement properties have been tested empirically in independent samples.

9. Practical Implications for Higher Education

9.1 Teach Verification as a Skill

Universities should treat verification as a teachable academic practice rather than as an assumed student behaviour. Instruction can include exercises requiring students to:

- identify factual claims in AI-generated responses
- locate the original source of a claim
- compare AI-generated statements with peer-reviewed literature
- identify fabricated or unsupported references
- distinguish evidence from interpretation
- document uncertainty and limitations

9.2 Use Performance-Based Assessment

AI literacy should increasingly be assessed through what students can do, rather than solely through self-reported confidence.

For example, students can be given an AI-generated response containing a mixture of accurate, incomplete and unsupported claims and asked to:

- identify the claims;
- verify selected claims;
- classify the evidence;
- explain discrepancies;
- revise the response; and
- justify their final academic judgement.

Such tasks would provide a more direct assessment of verification competence.

9.3 Develop Risk-Sensitive AI Policies

Institutional policies can distinguish between low-risk and high-risk uses of generative AI. For low-risk activities such as brainstorming, verification requirements may be relatively limited. For high-risk academic claims, stronger source verification and human review should be expected. This risk-sensitive approach is more practical than assuming that every AI interaction requires exactly the same level of scrutiny.

9.4 Support Faculty Development

Faculty members also need guidance on designing assignments in which AI can be used productively without displacing students' responsibility for evidence evaluation.

Professional development can therefore focus on:

- designing AI-aware assessments
- teaching source verification
- recognising hallucinated or unsupported claims
- developing transparent AI-use requirements
- assessing students' reasoning rather than merely the final AI-assisted product

10. Limitations

Several limitations should be considered when interpreting the findings. First, the sample size is relatively small ($N = 104$) and the source dataset was collected using purposive sampling. The findings should therefore not be generalised to university students as a population without further evidence.

Second, the data are cross-sectional and based on self-reported survey responses. The study therefore cannot establish temporal or causal relationships.

Third, verification behaviour was measured through survey items rather than direct observation of students performing verification tasks. Self-reported verification may differ from actual verification performance.

Fourth, the instrument was not developed specifically as a dedicated measure of objective academic-information verification performance. Consequently, the present findings should not be interpreted as a direct measure of students' actual verification competence or scientific reasoning ability.

Fifth, although internal consistency and factorability diagnostics were satisfactory, the present study does not conduct confirmatory factor analysis, measurement invariance testing or comprehensive psychometric validation. The sample size also limits the extent to which complex latent-variable models could be justified.

Sixth, the trust coefficient in the multivariable model is relatively small and statistically close to the conventional significance threshold. It should therefore be interpreted cautiously.

Seventh, the usage-frequency variable is represented by numerical categories in the source dataset. Because the present analysis retains the source coding, the substantive interpretation of these categories depends on the original coding scheme.

Finally, because this is a secondary-data study, the present authors did not control the original sampling design, questionnaire administration or data-collection context. The source dataset should therefore be treated as the empirical foundation of the analysis while the present paper's contribution lies in its independent analytical and theoretical treatment.

11. Future Research

Future research should extend this work in several directions. First, researchers should develop and validate a performance-based instrument for AI-mediated academic information verification. Such an instrument could assess students' ability to identify claims, locate evidence, evaluate sources, recognise unsupported or fabricated information and make justified judgements.

Second, larger and more diverse samples should be used to examine whether the relationships observed here are replicated across universities, disciplines, academic levels and cultural contexts.

Third, future measurement studies should use exploratory and confirmatory factor analysis where sample sizes permit, followed by measurement-invariance testing across relevant groups.

Fourth, longitudinal and experimental designs are needed to determine whether targeted AI-literacy interventions actually improve verification performance.

Fifth, future models should examine the relationship between trust, reliance, perceived AI competence, verification behaviour and decision quality. In particular, the concept of calibrated trust warrants direct empirical testing.

Sixth, future studies should explicitly compare self-reported verification behaviour with observed verification performance. This would help determine whether reported verification translates into actual evidence-checking behaviour.

Finally, the proposed HEAVEN framework should be operationalised into measurable dimensions and evaluated using independent samples and objective performance-based tasks. Such research would determine whether the framework has predictive and educational utility beyond the present secondary dataset.

12. Conclusion

This study presents an independent secondary-data statistical analysis of a publicly available dataset containing responses from 104 university students who used AI chatbots for academic activities. The analysis indicates that AI-related digital literacy has a positive association with verification behaviour and the largest standardized coefficient in the multivariable model. In contrast, perceived ease of use and usage frequency were not statistically significant in the model. Trust showed a small negative adjusted association with verification behaviour, although the corresponding bivariate association was not statistically significant. This finding should be interpreted cautiously and should not be treated as evidence of a causal effect.

The overall pattern supports a distinction between AI-use fluency and verification competence. Students may find AI systems easy to use and may use them frequently without necessarily engaging in systematic information verification.

On this basis, the study proposes the HEAVEN framework—Human-centred Evaluation of AI-generated academic information through Verification, Evidence and Nuanced judgement. The framework places human judgement, evidence evaluation, verification and contextual reasoning at the centre of responsible academic engagement with generative AI.

The principal implication is therefore that higher education should strengthen the competencies required to evaluate AI-generated information appropriately, rather than treating AI adoption, frequency of use or technological ease as sufficient indicators of responsible AI engagement.

The present study provides an analytical and conceptual foundation for this approach. Future research should test the HEAVEN framework with larger and more diverse samples, validated measures and direct performance-based assessments.

References

- [1]. Association of College and Research Libraries. (2016). Framework for information literacy for higher education. American Library Association.
- [2]. Autio, C., Schwartz, R., Dunietz, J., Jain, A., Mouchel, P., Sekar, R. & Zingales, L. (2024). Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile (NIST AI 600-1). National Institute of Standards and Technology.
- [3]. Facione, P. A. (1990). Critical thinking: A statement of expert consensus for purposes of educational assessment and instruction. The California Academic Press.
- [4]. Farras, A. J., Lusiawati, D. A. & Simatupang, S. S. S. (2026). Survey Dataset for Exploring Students' Trust and Verification Behaviour in AI Chatbot Use [Data set]. Zenodo. doi:10.5281/zenodo.21135117.
- [5]. Kaur, K. & Pundir, P. S. (2026). Responsible AI readiness in higher education: A multilevel framework integrating faculty capability, professional judgement and institutional conditions. *International Journal for Multidisciplinary Research*, 8(5), 1-16. doi:10.36948/ijfmr.2026.v08i05.87389.
- [6]. Kasneci, E., Sessler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., Gasser, U., Groh, G., Günnemann, S., Hüllermeier, E., Krusche, S., Kutyniok, G., Michaeli, T., Nerdel, C., Pfeiffer, F., Poquet, O., Sailer, M., Schmidt, A., Seidel, T., Stadler, M. & Kasneci, G. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103, 102274. doi:10.1016/j.lindif.2023.102274.

- [7]. Lo, C. K. (2023). What is the impact of ChatGPT on education? A rapid review of the literature. *Education Sciences*, 13(4), 410.
- [8]. Miao, F. & Holmes, W. (2023). Guidance for generative AI in education and research. UNESCO.
- [9]. Zawacki-Richter, O., Marín, V. I., Bond, M. & Gouverneur, F. (2019). Systematic review of research on artificial intelligence applications in higher education—Where are the educators? *International Journal of Educational Technology in Higher Education*, 16, 39.
- [10]. Zhai, X., Chu, X., Chai, C. S., Jong, M. S.-Y., Istenic, A., Spector, M., Liu, J.-B., Yuan, J. & Li, Y. (2024). A review of artificial intelligence (AI) in education from 2010 to 2020. *Smart Learning Environments*, 11, 28.

Cite this Article:

Kaur, K., & Pundir, P. S. (2026). *Critical evaluation of AI-generated academic information in higher education: Empirical evidence on trust, AI-related digital literacy and verification behaviour. International Journal of Scientific Research in Modern Science and Technology (IJSRMST)*, 5(9), 1–12.

Journal URL: <https://ijrmst.com/> **DOI:** <https://doi.org/10.59828/ijrmst.v5i9.459>.



This work is licensed under a [Creative Commons Attribution-Non-Commercial 4.0 International License](https://creativecommons.org/licenses/by-nc/4.0/).

© The Author(s) 2026. IJSRMST Published by Surya Multidisciplinary Publication.

Disclaimer/Publisher's Note: *The author(s) are solely responsible for the content and claims of this article. The publisher and editors assume no legal liability for any errors, copyright violations, or damages arising from its use.*