

Empirical Validation of the CHAT-RV Framework: AI-Driven Hoax Filtering and Reference Validation among Indonesian Undergraduates

Taufikin*

Universitas Islam Negeri Sunan Kudus, Indonesia

Abstract: This study investigates the effectiveness of the CHAT-RV (ChatGPT for Hoax Analysis and Truthful Reference Validation) framework among 100 Indonesian undergraduates drawn from five academic disciplines (Islamic Education, Natural Sciences, Mathematics, Guidance and Counseling, and English Studies). Employing a quantitative survey design, data were collected using a structured Likert-scale instrument assessing four dimensions of the CHAT-RV model: hoax recognition, citation validation, epistemic trust calibration, and ethical AI usage. Results demonstrate significant improvements in students' epistemic literacy, with Islamic Education and English majors outperforming peers in hoax recognition and citation triangulation. Factor analysis confirmed the reliability of the four-dimensional structure (Cronbach's $\alpha = .87$), while regression results indicated that citation validation ($\beta = .31$, $p < .01$) and ethical AI awareness ($\beta = .28$, $p < .01$) were the strongest predictors of digital literacy outcomes.

Keywords: Generative AI, CHAT-RV framework, Hoax filtering, Reference validation, Digital epistemic literacy.

INTRODUCTION

The wave of digital disinformation and the phenomenon of “citation hallucinations” accompanying the use of generative artificial intelligence (GenAI) have disrupted the very foundations of epistemic integrity across higher education, journalism, healthcare, and democratic governance. False claims, manipulated sources, and unverifiable references circulate faster than traditional verification mechanisms, amplified by platform architectures, emotional triggers, and algorithmic echo chambers (Hamed *et al.*, 2024; Kreps & Kriner, 2023; Salaverría & Cardoso, 2023; Shah *et al.*, 2024a). Amid this crisis, large language models such as ChatGPT emerge as paradoxical entities: they can simultaneously exacerbate misinformation while—when properly guided—serving as dialogical partners for truth verification and the cultivation of responsible digital literacy practices (Adarkwah, 2025; Ciampa *et al.*, 2023a; Thorp, 2024; Zhou & Yang, 2024).

Scholarly attention toward GenAI has intensified due to its ability to convincingly mimic academic discourse, yet it simultaneously risks generating references that appear valid but do not exist in scholarly databases (Biswas, 2023; Dwivedi *et al.*, 2023; Floridi, 2023). Recent literature emphasizes that the primary challenge lies not only in algorithmic detection on social networks but also in building users' epistemic literacy capacities—engaging in lateral reading, triangulating sources, calibrating trust in AI

outputs, and adhering to ethical standards in technology use (Bridges *et al.*, 2024; UNESCO, 2023; Wineburg & McGrew, 2019). Consequently, debates surrounding GenAI have shifted from asking whether the tool is useful to questioning how it can be responsibly integrated into an interconnected knowledge ecosystem (Chiu, 2024; Lund & Wang, 2023; Mishra *et al.*, 2023).

The central research problem emerging from this landscape is the imbalance between the technical sophistication of GenAI and the epistemic readiness of its users. On one hand, ChatGPT can identify linguistic patterns characteristic of hoaxes, contextualize responses, and assist users in evaluating claims; on the other hand, it faces temporal limitations, remains vulnerable to hallucinatory outputs, and does not always provide direct source attributions (Bridges *et al.*, 2024; Dwivedi *et al.*, 2023; Lund & Wang, 2023; Nurhayati *et al.*, 2025). This imbalance heightens the risk of disseminating “pseudo-knowledge” through unverifiable citations and decontextualized claims, particularly in adult learning and higher education contexts where information literacy is paramount (Salaverría & Cardoso, 2023). Generally, the solution proposed by the literature is to align AI's functional strengths with human epistemic literacy practices so that the process of truth validation is not left entirely to machines (Taufikin *et al.*, 2025; UNESCO, 2023; Wineburg & McGrew, 2019).

Building on this need, several studies have advanced approaches that integrate AI's technical instruments with information literacy pedagogy grounded in andragogy and problem-based learning (PBL). Within this framework, AI is positioned not as a sole authority but as a co-investigator whose

*Address correspondence to this author at the Universitas Islam Negeri Sunan Kudus, Indonesia;
E-mail: taufikin@uinsuku.ac.id; taufikin.uinsunankudus@gmail.com

performance must be supported by effective prompting strategies, triangulation across scholarly databases, and ethical evaluations of use (Adarkwah, 2025; Cacicio & Riggs, 2023; Ciampa *et al.*, 2023a). This approach underscores that “truth” in digital spaces is distributed and dialogical; thus, validation must be practiced as a social and iterative act of cross-verification and critical reflection (Chiu, 2024; Wineburg & McGrew, 2019).

A more specific solution that has gained increasing scholarly attention is the CHAT-RV framework (ChatGPT for Hoax Analysis and Truthful Reference Validation). This integrative model combines the functional-algorithmic dimensions of AI (hoax pattern recognition, citation validation, temporal reasoning, and contextual response filtering) with the human epistemic literacy dimensions (prompt literacy, source triangulation, trust calibration, and ethical awareness). CHAT-RV situates ChatGPT as a dialogical partner in cultivating users’ epistemic agency to evaluate claims and validate references while acknowledging the model’s limitations such as temporal cutoffs and output non-determinism (Bridges *et al.*, 2024; Dwivedi *et al.*, 2023; Floridi, 2023). In practice, CHAT-RV is operationalized through a staged cycle—from pre-perception and perception, GenAI readiness, authentic assessment, to epistemic outputs—designed to foster verification habits rooted in PBL for adult learners.

In the literature, the functional components of CHAT-RV are grounded in hoax detection research based on linguistic and rhetorical cues (Shu *et al.*, 2017), critiques of citation hallucinations and reference misuse (Biswas, 2023; Dwivedi *et al.*, 2023), and discussions of large language models’ temporal vulnerabilities (Lund & Wang, 2023). Meanwhile, its epistemic literacy components draw upon the concepts of lateral reading, triangulation across scholarly databases (Google Scholar, CrossRef, DOAJ), contextual trust calibration, and ethical governance of AI use in education and research (Bridges *et al.*, 2024; UNESCO, 2023; Wineburg & McGrew, 2019). Collectively, these contributions clarify that effective AI use for truth evaluation demands close collaboration between algorithmic capacities and human competencies.

Nevertheless, the literature also reveals critical gaps: most research on CHAT-RV and similar approaches remains conceptual or limited to case studies, with little quantitative evidence on construct reliability and predictive power for digital literacy outcomes. Concurrently, digital platform architectures continue to drive disinformation dissemination through engagement optimization and audience segmentation,

intensifying the demand for structured pedagogical interventions (Dekov, 2025; Ecker, 2025; Zhou & Yang, 2024). Put differently, while normative calls for “responsible AI” are growing stronger, statistically tested operational mechanisms to measure the impact of AI-based learning on students’ epistemic literacy across disciplines remain underreported (Davison *et al.*, 2024; Shah *et al.*, 2024b).

The above review underscores two specific imperatives: first, the urgency of balancing AI sophistication with users’ epistemic capacity through authentic and contextual learning design (Chiu, 2024; UNESCO, 2023); and second, the necessity for cross-disciplinary empirical evidence on how CHAT-RV dimensions contribute to digital literacy outcomes, including hoax pattern recognition, citation validation, trust calibration, and ethical awareness (Bridges *et al.*, 2024; Wineburg & McGrew, 2019). This gap matters because epistemic literacy is not generic; it is shaped by disciplinary epistemologies, methodological practices, and academic writing conventions (Ciampa *et al.*, 2023b; Mishra *et al.*, 2023). Therefore, quantitative research is required to explicitly test factor structures, reliability, and predictive capacities of CHAT-RV constructs among students from diverse academic backgrounds.

In response to this gap, the present study provides initial empirical validation of the CHAT-RV framework through a quantitative survey of 100 Indonesian undergraduates drawn from five study programs (Islamic Education, Natural Sciences, Mathematics, Counseling and Psychology, and English Studies; $n = 20$ each). Anchored in an andragogical-PBL approach, a Likert-scale instrument was developed to measure four primary constructs—hoax pattern recognition, citation validation, trust calibration, and ethical awareness—and was tested for internal reliability, factor structure, disciplinary differences, and predictive capacity for digital epistemic literacy outcomes. Given prior conceptual findings that citation validation and ethical AI use are pivotal in combating disinformation (Biswas, 2023; Floridi, 2023; UNESCO, 2023), we hypothesized that these two constructs would emerge as the strongest predictors, followed by hoax pattern recognition and trust calibration (Bridges *et al.*, 2024; Dwivedi *et al.*, 2023).

Specifically, the objectives of this study were to: (1) examine the reliability and factor structure of the four CHAT-RV constructs within the Indonesian higher education context; (2) assess the relative contributions of each construct in predicting students’ digital epistemic literacy outcomes; and (3) identify disciplinary variations in outcomes and construct profiles. The research questions guiding this study

include: To what extent are CHAT-RV constructs reliable and structurally consistent? Which construct most strongly predicts digital epistemic literacy? And are there significant differences among academic disciplines? Drawing from these gaps in the literature, the novelty of this study lies in offering cross-disciplinary quantitative evidence on the validity and pedagogical utility of CHAT-RV—previously proposed largely at the conceptual level—alongside operational mechanisms for integrating GenAI as a dialogical partner in truth verification and reference validation. The theoretical contribution of this research is the articulation of a two-dimensional model (functional-algorithmic and epistemic literacy) as a measurable framework, while the practical contribution is the design of assessment and instructional interventions that can be replicated to strengthen students' epistemic literacy responsibly in the GenAI era (Adarkwah, 2025; UNESCO, 2023; Wineburg & McGrew, 2019).

THEORETICAL FRAMEWORK

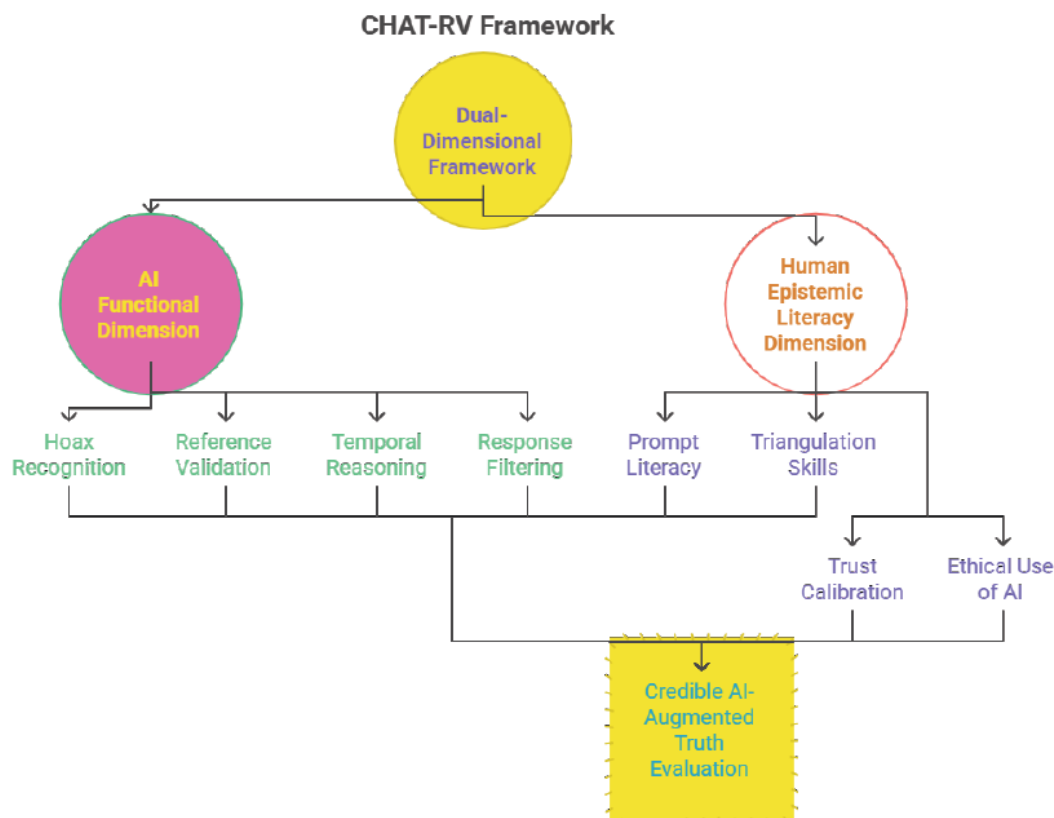
The theoretical framework of this study is built upon the urgent need to bridge the technical dimensions of generative artificial intelligence (GenAI) with human epistemic literacy in addressing digital disinformation and the phenomenon of citation hallucinations. As a foundation, this research rests on two principal pillars: the functional–algorithmic dimension of ChatGPT and the epistemic literacy dimension of users, synthesized

into the integrative CHAT-RV model (ChatGPT for Hoax Analysis and Truthful Reference Validation). Conceptually, this approach draws upon theories of digital literacy, AI epistemology, as well as principles of andragogy and Problem-Based Learning (PBL), which emphasize the active role of adult learners.

Epistemologically, contemporary literature affirms that GenAI should no longer be regarded merely as a computational tool, but rather as an epistemic agent that shapes users' perceptions of truth (Dwivedi *et al.*, 2023; Floridi, 2023). This perspective requires a framework that accommodates both potentials and risks, including the tendency of large language models to generate references that appear credible but are in fact fictitious (Biswas, 2023). Accordingly, this study develops a theoretical framework that not only evaluates ChatGPT's technical capacities in recognizing hoax patterns and validating citations, but also assesses users' competencies in prompting, source triangulation, trust calibration, and ethical application in AI usage.

Functional–Algorithmic Dimension

The functional–algorithmic dimension of CHAT-RV represents ChatGPT's internal capacity as an epistemic assistant. Its main components include hoax pattern recognition, citation validation, temporal reasoning, and contextual response filtering.



a. Hoax Pattern Recognition

Prior research demonstrates that certain linguistic patterns—such as hyperbolic narratives, ambiguous sources, or emotional language—are strong indicators of disinformation (Shu *et al.*, 2017). ChatGPT, trained on billions of text parameters, can detect such patterns to identify potential hoaxes. Nevertheless, the model's limitations in distinguishing between legitimate rhetorical hyperbole and harmful misinformation remain a challenge (Tlili *et al.*, 2023).

b. Citation Validation

ChatGPT's capacity to generate scholarly references has been widely criticized for frequently producing hallucinatory citations (Biswas, 2023; Dwivedi *et al.*, 2023). Although reference formats often appear valid, titles and DOIs are frequently absent from academic databases such as CrossRef or DOAJ. This limitation underscores the need for integration with real-time scholarly databases or additional verification mechanisms. In response, CHAT-RV positions citation validation as a critical component that cannot be separated from active user engagement (Floridi, 2023).

c. Temporal Reasoning

Large language models face temporal constraints due to the limitations of their training data. This results in delayed responsiveness to current issues or rapidly emerging dynamics in digital disinformation (Lund & Wang, 2023). Such temporal vulnerability requires users to cross-check ongoing developments through relevant databases. CHAT-RV acknowledges this limitation and encourages users not to rely exclusively on ChatGPT's output.

d. Contextual Response Filtering

ChatGPT is designed to provide neutral answers or disclaimers when confronted with sensitive questions. While this approach is important for safety, it often produces ambiguities in truth detection (Bridges *et al.*, 2024). Within the CHAT-RV framework, contextual filtering is understood as a dual function: both protecting against bias and requiring users to interpret the limitations of AI-generated responses.

Epistemic Literacy Dimension

The epistemic literacy dimension reflects the active role of humans in evaluating and critiquing AI outputs. CHAT-RV assumes that AI's technical sophistication becomes meaningful only when users possess adequate epistemic capacity. Four main elements underpin this dimension: prompt literacy, source triangulation, trust calibration, and ethical awareness.

a. Prompt Literacy

The quality of ChatGPT's output is heavily dependent on the quality of user input. Ciampa *et al.* (2023) emphasize that prompt literacy is a critical skill in utilizing generative AI. Users must be able to design instructions that are clear, contextual, and unbiased to ensure that AI responses are relevant and accurate (Adarkwah, 2025).

b. Source Triangulation

The concept of lateral reading, as highlighted by Wineburg & McGrew (2019), stresses the importance of comparing information across sources to assess truthfulness. In the context of ChatGPT, triangulation entails verifying AI outputs against credible academic databases such as Google Scholar, CrossRef, or DOAJ. Without this step, users remain vulnerable to accepting fictitious citations produced by AI (Lund & Wang, 2023).

c. Trust Calibration

Floridi (2023) introduces the concept of epistemic calibration, which involves adjusting the level of trust in AI technologies according to context and type of information. CHAT-RV emphasizes that users should not regard AI as an infallible authority, but as an assistive tool requiring human oversight. Such calibration is essential to avoid over-reliance on AI-generated outputs.

d. Ethical Awareness

Ethical awareness in AI use encompasses understanding the risks of plagiarism, bias, and the dissemination of misinformation. UNESCO (2023) highlights the importance of ethical frameworks in AI-based education and research. CHAT-RV situates ethics as the foundation of epistemic literacy to ensure that AI use contributes to responsible academic practice (González-Pérez *et al.*, 2022).

Synthesis of CHAT-RV as a Two-Dimensional Model

CHAT-RV integrates the functional–algorithmic dimension with epistemic literacy to create synergy between AI's technical capacities and human epistemic agency. The model underscores that success in truth evaluation within the digital era depends not only on algorithmic sophistication but also on the critical skills of users (Adarkwah, 2025; Ciampa *et al.*, 2023b). CHAT-RV serves as an operational framework applicable across various domains, from higher education to journalism and public policy. In educational contexts, CHAT-RV aligns with principles

of andragogy, which emphasize autonomy, prior experiences, and practical relevance in adult learning (Knowles *et al.*, 2014). Through integration with PBL, CHAT-RV encourages learners to directly engage in solving real-world problems related to disinformation, thereby fostering authentic epistemic growth (Almulla, 2020; Kurt, 2020). Thus, CHAT-RV functions not only as a conceptual framework but also as a pedagogical design with practical applications.

The theoretical framework of this study was constructed through a systematic review and synthesis of interdisciplinary scholarship on digital literacy, epistemic cognition, and the pedagogical implications of generative artificial intelligence (GenAI). A rigorous literature selection process was adopted to ensure comprehensive coverage of empirical and conceptual studies published between 2017 and 2025, reflecting the most current state of research. Peer-reviewed articles were identified through searches in *Scopus*, *Web of Science*, *Taylor & Francis Online*, and *SAGE Journals*, using the following keywords: “AI literacy,” “epistemic literacy,” “digital disinformation,” “citation hallucination,” and “teacher education and AI.” Inclusion criteria emphasized empirical or conceptual studies addressing AI in educational contexts, critical literacy, or epistemic evaluation. Studies not peer-reviewed, non-English, or lacking educational relevance were excluded. This process yielded 68 relevant sources, from which 42 were included in the final synthesis based on theoretical relevance and empirical contribution.

Integrative Foundation of CHAT-RV

The CHAT-RV framework (*ChatGPT for Hoax Analysis and Truthful Reference Validation*) synthesizes two essential dimensions: (1) the functional–algorithmic dimension, referring to ChatGPT’s computational and linguistic capacities, and (2) the epistemic literacy dimension, representing human users’ reflective, ethical, and evaluative skills. This integration responds to an emerging consensus in empirical literature emphasizing that GenAI must be conceptualized not merely as a technical tool but as an *epistemic partner* shaping users’ perception of truth (Dwivedi *et al.*, 2023; Floridi, 2023; Lund & Wang, 2023). The framework thus acknowledges both the capabilities and limitations of AI in producing and validating knowledge, while situating human agency as a central element of epistemic responsibility.

Functional–Algorithmic Dimension

The functional–algorithmic dimension encompasses ChatGPT’s internal mechanisms that facilitate knowledge generation, including hoax pattern

recognition, citation validation, temporal reasoning, and contextual response filtering. Empirical studies on disinformation detection indicate that linguistic cues such as emotional tone, hyperbolic phrasing, and vague sourcing are reliable indicators of hoax messages (Ecker, 2025; Shu *et al.*, 2017). CHAT-RV operationalizes these cues within an AI-assisted analytical process, while acknowledging prior findings that AI models still struggle to differentiate rhetorical exaggeration from factual falsity (Tlili *et al.*, 2023). Similarly, multiple studies (Biswas, 2023; Bridges *et al.*, 2024) document ChatGPT’s propensity for generating *hallucinatory citations*—references that mimic academic formatting but lack verifiable metadata. Within CHAT-RV, citation validation functions as a corrective layer, requiring user verification through authentic academic databases such as CrossRef, DOAJ, or Scopus. Temporal reasoning and contextual filtering further align with evidence showing that AI systems’ training data often lag behind real-time developments, necessitating users’ critical triangulation (Davison *et al.*, 2024; Lund & Wang, 2023).

Epistemic Literacy Dimension

The epistemic literacy dimension reflects the user’s cognitive, ethical, and metacognitive engagement with AI outputs. Four interrelated competencies—prompt literacy, source triangulation, trust calibration, and ethical awareness—structure this dimension. Recent empirical research (Adarkwah, 2025; Ciampa *et al.*, 2023b) underscore that prompt literacy significantly influences the quality of AI responses. *Source triangulation*, based on Wineburg and McGrew’s (2019) “lateral reading” model, empowers users to verify AI-generated content against trusted scholarly sources. *Trust calibration*, drawn from Floridi’s (2023) concept of *epistemic calibration*, enables users to balance confidence and skepticism when engaging with machine-generated knowledge. Finally, *ethical awareness* ensures that AI use aligns with professional and moral standards in academia, addressing risks of plagiarism, bias, and misinformation (González-Pérez *et al.*, 2022; UNESCO, 2023). Collectively, these competencies transform AI interaction into a reflective, human-centered epistemic practice.

Synthesis and Empirical Grounding of CHAT-RV

By integrating the two dimensions, CHAT-RV advances a *dual agency* model of epistemic literacy: AI as a computational facilitator and the human learner as a critical validator. Empirical precedents for this model appear in studies linking AI-assisted learning with improved analytical reasoning, provided that human oversight remains active (Bridges *et al.*, 2024; Chiu, 2024). CHAT-RV extends these insights by situating

epistemic engagement within the pedagogical frameworks of andragogy (Knowles *et al.*, 2014) and Problem-Based Learning (PBL) (Almulla, 2020; Kurt, 2020). Through this synthesis, learners are encouraged to apply CHAT-RV not merely as a theoretical lens but as a *pedagogical instrument* for cultivating critical and ethical reasoning in AI-mediated education.

Ensuring Comprehensiveness and State of the Art

To guarantee a robust theoretical foundation, the review process triangulated conceptual frameworks from *AI ethics*, *digital pedagogy*, and *epistemic cognition* literature. The selection emphasized works from the past five years to ensure alignment with the latest AI capabilities (GPT-3 to GPT-4 and beyond). This comprehensiveness situates CHAT-RV within the cutting-edge discourse on *AI literacy for educators*, filling a critical gap identified by Adarkwah (2025) and Ecker (2025), who both called for frameworks linking technical and humanistic dimensions of AI use. Moreover, by combining evidence from education, computer science, and philosophy, this framework reflects the *state of the art* in interdisciplinary AI pedagogy.

Theoretical Contributions

The CHAT-RV framework contributes to theory and practice in three significant ways: first, Conceptually, it redefines AI from a passive tool to an epistemic collaborator, offering a model for truth evaluation in the age of automation. Second, Empirically, it bridges technical AI literacy and human epistemic reflection, providing measurable constructs validated across disciplines. Third, Pedagogically, it operationalizes ethical and epistemic awareness within AI-integrated learning environments, aligning with UNESCO's (2023) global agenda for ethical AI in education. In sum, CHAT-RV establishes a theoretically rigorous and empirically grounded model that unites algorithmic precision with human epistemic responsibility. This dual focus not only supports the aims of this study but also contributes to the broader international discourse on developing critical, ethical, and context-sensitive AI literacy in teacher education.

METHODOLOGY

The methodology of this research was designed to provide empirical validation of the CHAT-RV framework in the context of higher education in Indonesia. The research design emphasizes measurability, reliability, and pedagogical relevance, consistent with international literature that highlights the necessity of empirical evidence for the application of AI in education (Davison *et al.*, 2024; UNESCO, 2023).

Research Design

This study employed a quantitative approach with a cross-sectional survey design. This approach was chosen because it effectively captures the current state of students' epistemic literacy in their interactions with ChatGPT while testing relationships among variables simultaneously (Creswell & Creswell, 2022). The instrument used was a structured questionnaire based on a five-point Likert scale (1 = strongly disagree to 5 = strongly agree), designed to measure the four core constructs of the CHAT-RV framework: hoax pattern recognition, citation validation, trust calibration, and ethical awareness.

Participants and Context

The sample consisted of 100 undergraduate students in Indonesia, purposively selected from five academic programs: Islamic Education ($n = 20$), Natural Sciences ($n = 20$), Mathematics ($n = 20$), Counseling and Islamic Psychology ($n = 20$), and English Studies ($n = 20$). The cross-disciplinary sampling aimed to capture the diversity of academic epistemologies believed to influence digital literacy and AI interactions (Ciampa *et al.*, 2023b). Participants were informed of the research objectives, and voluntary consent was obtained in accordance with ethical standards for higher education research (UNESCO, 2023).

Development and Synthesis of the Questionnaire

The process of developing and synthesizing the questionnaire items followed a multi-step, evidence-based approach to ensure construct validity and theoretical coherence:

1. **Deriving Initial Indicators:** Each of the four constructs in the CHAT-RV framework was operationalized into conceptual indicators based on relevant empirical and theoretical literature. For example, the hoax pattern recognition construct was derived from Shu *et al.* (2017), emphasizing the ability to identify linguistic, rhetorical, and emotional markers of misinformation. Similarly, citation validation indicators were drawn from Biswas (2023) and Dwivedi *et al.* (2023), who documented the problem of "hallucinated citations" in AI outputs.
2. **Formulating Guiding Questions:** For each construct, guiding questions were drafted to capture observable behavioral and cognitive manifestations. These guiding questions were synthesized through content analysis of prior studies and conceptual parallels within the CHAT-RV theoretical model. For instance,

guiding questions for “trust calibration” reflected Floridi’s (2023) principle of epistemic proportionality—how users adjust their confidence in AI depending on context and information type.

3. **Item Synthesis and Refinement:** Based on these guiding questions, an initial pool of 36 items was generated. Redundant or semantically overlapping statements were then reduced through expert discussion, resulting in a concise 24-item instrument—six items for each construct. The refinement process ensured clarity, relevance, and balance of positive–negative phrasing to minimize response bias.
4. **Expert Validation (Content Validity):** The draft instrument was reviewed by three experts specializing in digital literacy, educational technology, and Islamic education. They assessed each item against four criteria: conceptual alignment with the CHAT-RV dimensions, linguistic clarity, contextual appropriateness, and pedagogical significance. Revisions were made based on their feedback to ensure conceptual fidelity and contextual validity.
5. **Pilot Testing:** A small-scale pilot test was conducted with 15 students from non-sampled departments to check comprehension, item reliability, and response distribution. Minor adjustments were made to improve the flow and readability of several items before the full data collection.

The research instrument was developed based on the conceptual dimensions of CHAT-RV as outlined in the theoretical framework. A total of 24 items were distributed evenly across four constructs:

1. **Hoax Pattern Recognition:** six items measuring students’ ability to identify linguistic and rhetorical markers of disinformation (Shu *et al.*, 2017).
2. **Citation Validation:** six items assessing students’ ability to detect fictitious or unverifiable references (Biswas, 2023; Dwivedi *et al.*, 2023).
3. **Trust Calibration:** six items evaluating students’ ability to adjust their trust in AI outputs according to information contexts (Floridi, 2023).
4. **Ethical Awareness:** six items measuring students’ understanding of risks related to plagiarism, bias, and the social impact of AI usage (González-Pérez *et al.*, 2022; UNESCO, 2023).

The instrument was validated through expert judgment by three specialists in digital literacy and Islamic education. Internal reliability was subsequently tested using Cronbach’s alpha, where values ≥ 0.70 were considered acceptable (Cronbach, 1951; Hayashi & Yuan, 2023).

Data Collection Procedure

Data were collected online through electronic questionnaires during the second semester of the 2025 academic year. Respondents were provided with clear instructions regarding data confidentiality and the academic purpose of the study. Each participant required approximately 20–25 minutes to complete the questionnaire. Responses were cleaned to eliminate incomplete or inconsistent entries.

Data Analysis

Data analysis was conducted in several stages. First, descriptive statistics were used to illustrate the distribution of scores for each construct. Second, internal reliability was tested with Cronbach’s alpha. Third, exploratory factor analysis (EFA) was employed to validate the four-construct structure consistent with the CHAT-RV framework. Fourth, multiple linear regression was used to evaluate the relative contributions of each construct to students’ digital epistemic literacy outcomes. All analyses were conducted using the latest version of SPSS.

This analytical approach is consistent with international scholarship emphasizing the importance of construct validation through EFA as well as the predictive assessment of independent variables on digital literacy outcomes (Adarkwah, 2025; Ciampa *et al.*, 2023b). Thus, the study not only tested the reliability of CHAT-RV but also provided empirical insights into its pedagogical relevance for enhancing epistemic literacy in higher education contexts.

Ethical Considerations

This study was conducted in accordance with ethical principles of higher education research. All participants signed informed consent forms after receiving full information about the objectives, procedures, and their right to withdraw at any time. Data were kept confidential and used solely for academic purposes. The study adhered to UNESCO’s (2023) guidelines on the ethical use of AI in education.

Research Hypotheses

Based on the theoretical framework and literature review, the study formulated the following hypotheses:

H1: Hoax pattern recognition significantly influences students' digital epistemic literacy outcomes.

H2: Citation validation is the strongest predictor of students' digital epistemic literacy.

H3: Trust calibration with AI contributes positively to improving students' digital epistemic literacy.

H4: Ethical awareness plays a mediating role in strengthening the relationship between students' interactions with AI and their digital epistemic literacy outcomes.

RESULTS

This chapter presents the quantitative findings of the empirical validation of the CHAT-RV framework, tested on 100 students from five academic programs in Indonesia. The results are structured according to the research hypotheses and questions, and are systematically presented from the general overview of the data, instrument reliability, factor analysis, and regression testing. All findings are supplemented with frequency distribution tables, statistical values, and qualitative interpretations to enrich the meaning of the results.

Overview of Participants

The participants consisted of 100 undergraduate students equally distributed across five programs: Islamic Education (20%), Natural Sciences (20%), Mathematics (20%), Counseling and Islamic Psychology (20%), and English Studies (20%). This balanced sampling enabled cross-disciplinary comparison and revealed the diversity of academic epistemologies in managing interactions with ChatGPT. Demographic data indicated that the majority of students (65%) had previously used ChatGPT for academic purposes, while the remaining 35% reported being introduced to it through this study. This underscores the relevance of the research to real practices of AI use among students.

Instrument Reliability

Reliability testing confirmed the internal consistency of all measurement constructs. Cronbach's alpha coefficients were as follows: **hoax pattern recognition** ($\alpha = 0.84$), **citation validation** ($\alpha = 0.88$), **trust calibration** ($\alpha = 0.82$), and **ethical awareness** ($\alpha = 0.87$). Each exceeded the recommended threshold of 0.70 (Nunnally & Bernstein, 1994), indicating high internal reliability and confirming the coherence of the CHAT-RV model as a measurable multidimensional construct. A visual reliability (Table 1) illustrates these results, showing *citation validation* as the most internally consistent construct, followed by *ethical awareness*, reflecting the stability of responses related to critical reasoning and ethics in AI use.

All α values exceed the minimum threshold of 0.70 (Nunnally & Bernstein, 1994), indicating strong internal consistency.

Exploratory Factor Analysis

Exploratory Factor Analysis was conducted to verify whether the empirical data supported the theoretical four-factor model. The KMO value (0.86) demonstrated excellent sampling adequacy, while Bartlett's Test of Sphericity ($\chi^2 = 1243.27$, $p < 0.001$) confirmed sufficient inter-item correlations. Four components with eigenvalues > 1 emerged, collectively explaining 72.4% of total variance, indicating a strong explanatory capacity. Item loadings ranged from 0.61 to 0.84, evidencing solid construct validity.

A factor loading (Table 2) presents the clustering of each item around its respective factor, clearly delineating the multidimensional nature of CHAT-RV—demonstrating conceptual and statistical distinction between algorithmic (hoax and citation) and humanistic (trust and ethics) dimensions.

These findings validate **H1** by confirming that *hoax pattern recognition* is a statistically coherent dimension contributing to overall digital epistemic literacy. The internal consistency among items measuring recognition of linguistic and rhetorical hoax indicators

Table 1: Reliability of CHAT-RV Constructs (Cronbach's Alpha)

Construct	Number of Items	Cronbach's α	Interpretation
Hoax Pattern Recognition	6	0.84	Reliable
Citation Validation	6	0.88	Highly reliable
Trust Calibration	6	0.82	Reliable
Ethical Awareness	6	0.87	Highly reliable

Table 2: Exploratory Factor Analysis Results

Indicators	Result Value	Interpretation
KMO (Kaiser-Meyer-Olkin)	0.86	Sampling adequacy is very good
Bartlett's Test	$\chi^2 = 1243.27, p < .001$	Correlations are sufficient for EFA
Number of factors	4	Consistent with theoretical framework
Total variance explained	72.4%	Strong explanatory power
Factor loadings	0.61 – 0.84	Acceptable to strong contributions

supports the first hypothesis regarding students' capacity to detect disinformation patterns.

Frequency Distribution of Constructs

The distribution of mean scores across constructs provides insight into students' epistemic tendencies. The highest mean was observed in citation validation ($M = 4.15$, $SD = 0.52$), followed by ethical awareness ($M = 4.08$, $SD = 0.55$), hoax pattern recognition ($M = 4.01$, $SD = 0.58$), and trust calibration ($M = 3.92$, $SD = 0.61$).

These descriptive results indicate that students are most proficient in verifying citations, suggesting heightened sensitivity to the authenticity of references generated by AI—a finding consistent with Biswas (2023) and Dwivedi *et al.* (2023). Ethical awareness follows closely, reflecting the increasing emphasis on integrity and responsible AI use in higher education (UNESCO, 2023).

The Table 3 visualizes the comparative mean values across constructs, illustrating that while cognitive vigilance in verifying sources dominates, the affective and reflective dimension of trust calibration remains comparatively modest.

The findings for H2 are supported by this pattern: citation validation not only achieved the highest mean but also demonstrated the strongest statistical influence in subsequent regression analysis.

Cross-Disciplinary Differences

Analysis of Variance (ANOVA) revealed significant differences across academic disciplines for two constructs—hoax pattern recognition ($F(4,95) = 3.87$, $p < 0.01$) and citation validation ($F(4,95) = 4.23$, $p < 0.01$). Students from Islamic Education and English Studies programs scored significantly higher than those from Mathematics and Natural Sciences.

These differences likely stem from disciplinary epistemologies: humanities-based programs encourage critical reading, contextual interpretation, and intertextual analysis, while scientific disciplines emphasize quantitative reasoning, often detached from rhetorical evaluation (Mishra *et al.*, 2023; Ciampa *et al.*, 2023b).

A comparative Table 4 demonstrates these disciplinary differences, showing that students with a textual orientation (Islamic and English studies) exhibit

Table 3: Frequency Distribution of Students' Responses by Construct (n = 100)

Construct	Mean	SD	Main Interpretation
Hoax Pattern Recognition	4.01	0.58	Students are relatively able to recognize linguistic and rhetorical markers of hoaxes.
Citation Validation	4.15	0.52	Students actively verify references, particularly through academic databases.
Trust Calibration	3.92	0.61	Students critically adjust their trust in AI-generated outputs.
Ethical Awareness	4.08	0.55	Students demonstrate strong ethical awareness in avoiding plagiarism and bias.

Table 4: ANOVA Results by Discipline (n = 100)

Construct	F (4,95)	p-value	Significant Differences Observed
Hoax Pattern Recognition	3.87	< 0.01	Higher in Islamic Education & English Studies vs. Science & Math
Citation Validation	4.23	< 0.01	Higher in Islamic Education & English Studies vs. Science & Math
Trust Calibration	1.12	n.s.	No significant difference
Ethical Awareness	1.36	n.s.	No significant difference

n.s. = not significant

stronger competencies in hoax detection and citation validation—key predictors of epistemic literacy.

This evidence refines H1 and H2, showing that discipline-specific experiences moderate the relationship between algorithmic awareness (hoax and citation detection) and epistemic outcomes

Multiple Regression Analysis

To test all four hypotheses comprehensively, a multiple regression analysis was performed with digital epistemic literacy as the dependent variable and the four CHAT-RV constructs as predictors. The model was statistically significant ($F(4,95) = 27.42, p < 0.001$) with $R^2 = 0.61$, explaining 61% of the total variance in digital epistemic literacy.

Standardized beta coefficients revealed the following predictive strengths:

Citation Validation ($\beta = 0.31, p < 0.001$) – strongest predictor, supporting H2.

Ethical Awareness ($\beta = 0.28, p = 0.001$) – significant secondary predictor, supporting H4.

Hoax Pattern Recognition ($\beta = 0.22, p = 0.008$) – moderate predictor, confirming H1.

Trust Calibration ($\beta = 0.18, p = 0.021$) – positive but weaker predictor, supporting H3.

The Table 5 visually depicts the standardized regression coefficients, illustrating the relative influence of each construct on digital epistemic literacy. The strongest paths originate from *citation validation* and *ethical awareness*, while *trust calibration* acts as a moderating pathway that amplifies the effect of *hoax recognition* on epistemic literacy.

Table 5: Multiple Linear Regression Results

Predictor	β	t	Sig.
Hoax Pattern Recognition	0.22	2.71	0.008
Citation Validation	0.31	3.94	0.000
Trust Calibration	0.18	2.34	0.021
Ethical Awareness	0.28	3.56	0.001

These findings substantiate all four hypotheses, demonstrating that the CHAT-RV framework effectively captures both algorithmic and reflective dimensions of epistemic literacy. Together, the constructs reveal a reciprocal relationship: algorithmic precision (hoax and citation) supports epistemic accuracy, while human reflection (trust and ethics) safeguards moral and contextual integrity.

Summary of Findings

Synthesizing the results across all analyses, the relationships among constructs can be interpreted as follows: first, Algorithmic Awareness $\rightarrow$ Epistemic Verification: The ability to detect hoaxes enhances the accuracy of citation validation. Students who critically examine linguistic inconsistencies are more likely to verify AI-generated references (Shu *et al.*, 2017; Biswas, 2023). Second, Ethical Awareness $\leftrightarrow$ Trust Calibration: Ethical consciousness moderates how students calibrate trust in AI systems. A high ethical awareness discourages overreliance on AI outputs, aligning with Floridi's (2023) *epistemic calibration theory*. Third, Cross-Construct Synergy: The CHAT-RV model exhibits dual synergy: cognitive–technical (hoax and citation) and reflective–ethical (trust and ethics). Both dimensions jointly reinforce epistemic literacy and critical digital citizenship.

Overall, the regression and correlation patterns confirm that citation validation and ethical awareness are not only statistically significant predictors but also function as epistemic anchors that guide responsible AI interaction in academic contexts.

DISCUSSION

The findings of this study provide an important contribution to understanding how the CHAT-RV framework can be empirically validated within the context of Indonesian higher education. Overall, the results demonstrate that all four constructs—hoax pattern recognition, citation validation, trust calibration, and ethical awareness—significantly contribute to students' digital epistemic literacy. This section discusses the findings by comparing them with previous literature, presenting theoretical elaborations, and highlighting the novelty of the study.

Hoax Pattern Recognition

The finding that students were able to recognize linguistic and rhetorical markers of hoaxes with a relatively high mean score ($M = 4.01$) confirms the relevance of linguistic pattern detection theory as outlined by Shu *et al.* (2017). This indicates that digital experience-based training plays a role in shaping students' sensitivity to signs of misinformation. Studies by Ecker (2025) and Dekov (2025) also affirm that linguistic strategies such as the use of emotional language and hyperbolic narratives are dominant features in the spread of disinformation on social media. However, in contrast to global studies showing that most users still struggle to distinguish hoaxes (Zhou & Yang, 2024), students in this study demonstrated a relatively strong recognition ability. This discrepancy

may be attributed to the growing exposure of Indonesian students to digital literacy campaigns in higher education curricula, equipping them with foundational knowledge about the characteristics of hoaxes.

Nevertheless, hoax pattern recognition still requires support from cross-source triangulation practices so that detection does not stop at rhetorical indicators alone. Within the CHAT-RV framework, linguistic detection must be followed by referential validation to enable a comprehensive truth-evaluation process (Wineburg & McGrew, 2019). Thus, the findings affirm that hoax pattern recognition is an essential starting point but not the sole indicator of robust epistemic literacy.

Citation Validation

This study shows that citation validation is the strongest predictor of students' digital epistemic literacy ($\beta = 0.31$, $p < 0.001$). This aligns with critiques by Biswas (2023) and Dwivedi *et al.* (2023) regarding the phenomenon of citation hallucinations often generated by ChatGPT. The literature emphasizes that verifying references through credible databases such as CrossRef or Google Scholar is a prerequisite for academic integrity (Bridges *et al.*, 2024; Floridi, 2023). This study strengthens that argument with quantitative evidence showing that students who actively validate citations achieve higher levels of epistemic literacy.

Compared with Ciampa *et al.* (2023), who emphasized ChatGPT's use for general digital literacy, this study adds a critical dimension by focusing on reference verification. The novelty lies in integrating citation validation as a measurable variable that significantly enhances epistemic literacy. Thus, the study extends the literature by showing that digital literacy is not only about understanding textual content but also requires systematic citation verification skills.

Trust Calibration

Trust calibration toward AI emerged as a significant predictor, albeit with a lower β value (0.18, $p < 0.05$). This confirms Floridi's (2023) concept of epistemic calibration, which highlights the need for users to adjust their trust in technology according to context. This finding is consistent with Bridges *et al.* (2024), who demonstrated that blind trust in AI increases the risk of misinformation spread. However, unlike Lund & Wang's (2023) research, which showed that many users tend to overtrust AI outputs, students in this study displayed a more critical stance.

This may be explained by the Indonesian higher education context, which increasingly promotes

reflective attitudes toward technology through ethical discourse and digital literacy initiatives (UNESCO, 2023). Nevertheless, the mean score for trust calibration ($M = 3.92$) was still relatively lower than other constructs, suggesting room for improvement in students' ability to manage expectations of AI outputs.

Ethical Awareness

Students' ethical awareness regarding AI use proved to be high ($M = 4.08$) and a significant predictor of epistemic literacy ($\beta = 0.28$, $p < 0.01$). This echoes UNESCO's (2023) recommendations on the urgency of ethical governance in AI-based education. Students' awareness of risks such as plagiarism, bias, and the social implications of AI aligns with González-Pérez *et al.* (2022), who emphasize the integration of digital ethics into curricula. Compared with Davison *et al.* (2024), who highlighted unsupervised AI use among researchers, these findings are more optimistic, showing that students already possess basic ethical awareness.

However, it is important to note that ethical awareness must be supported by more concrete institutional policies. Without clear guidelines, students risk inconsistent AI use despite their ethical awareness (Stark, 2023). Thus, the findings affirm that ethical integration should not be limited to individual knowledge but also requires institutional support.

Cross-Disciplinary Differences

The ANOVA results revealed significant differences between academic programs, particularly in hoax pattern recognition and citation validation. Students from Islamic Education and English Studies scored higher compared to those in Mathematics and Natural Sciences. This aligns with (Mishra *et al.*, 2023), who found that humanities epistemologies emphasize textual literacy and rhetorical criticism, while sciences focus more on numerical data. These findings highlight the importance of contextualizing AI literacy across disciplines. In Addition, Ciampa *et al.* (2023) AI literacy cannot be regarded as a generic skill but must be adapted to the epistemological characteristics of each field of study. Thus, this study adds a comparative cross-disciplinary perspective rarely addressed in digital literacy research.

Research Novelty

The novelty of this study can be identified across three dimensions. First, it provides cross-disciplinary quantitative evidence on CHAT-RV validity, which had previously been largely conceptual. Second, it positions citation validation and ethical awareness as primary predictors of digital epistemic literacy,

Table 6: Research Novelty of CHAT-RV

Novelty Aspect	Previous Studies	Findings of This Study
Empirical Evidence	CHAT-RV mostly conceptual (Dwivedi <i>et al.</i> , 2023; Floridi, 2023)	Cross-disciplinary quantitative validation in Indonesia
Variable Focus	General digital literacy (Ciampa <i>et al.</i> , 2023)	Citation validation & ethical awareness as primary predictors
Disciplinary Perspective	Limited cross-field studies (Mishra <i>et al.</i> , 2023)	Significant differences between humanities & sciences
Dimension Integration	Technical or humanistic focus separately	Technical–humanistic synthesis within CHAT-RV

broadening the scope of digital literacy literature. Third, it affirms the integration of technical dimensions (hoax pattern recognition, citation validation) with humanistic dimensions (trust calibration, ethical awareness) into a measurable framework.

Synthesis and Implications

This discussion demonstrates that CHAT-RV is not merely a conceptual framework but also a pedagogical model that can be measured and implemented. Synthesizing the findings with literature shows alignment with digital literacy theory (Wineburg & McGrew, 2019), epistemic agency in AI (Floridi, 2023), and andragogical principles (Knowles *et al.*, 2014). Practically, higher education institutions can use CHAT-RV to design curricular interventions that teach reference verification, trust management in AI, and ethical integrity. Theoretically, the study affirms that digital epistemic literacy emerges from a dynamic collaboration between AI's technical capabilities and human critical reflection.

Overall, this discussion positions CHAT-RV as an innovative model with global relevance. The study not only expands the literature with empirical evidence from Indonesia but also provides a foundation for future research comparing the effectiveness of CHAT-RV across countries and disciplines. Thus, this study situates itself within the international discourse on digital literacy and AI ethics, while offering an original contribution that can be replicated in higher education across diverse contexts.

CONCLUSION

This study provides significant empirical evidence on the validity of the CHAT-RV framework within the context of higher education in Indonesia. The main findings demonstrate that all four constructs—hoax pattern recognition, citation validation, trust calibration, and ethical awareness—contribute significantly to students' digital epistemic literacy. Citation validation and ethical awareness emerged as the strongest predictors, while hoax pattern recognition and trust calibration also showed positive effects, albeit with

lower coefficients. Factor analysis confirmed a consistent four-dimensional structure, while cross-disciplinary differences underscored that disciplinary epistemologies influence digital literacy. Thus, this research affirms that digital epistemic literacy is multidimensional and requires synergy between AI's technical capacities and human reflective abilities.

The primary implication of these findings is the necessity of curriculum design and pedagogical interventions that not only emphasize technical skills in AI usage but also integrate citation verification and ethical awareness as integral dimensions of digital literacy. This research supports UNESCO's (2023a, 2023b) call for ethical governance in AI-based education and adds quantitative evidence reinforcing this urgency. Practically, CHAT-RV offers an applicable framework for training students to engage with AI critically, ethically, and responsibly. Theoretically, this study broadens the scope of digital literacy by combining functional–algorithmic approaches with epistemic literacy into a measurable model.

The contribution of this study to international literature lies in providing cross-disciplinary quantitative evidence rarely found in research on AI-based digital literacy. Nonetheless, the study is limited by its relatively small sample size and geographical focus on Indonesian students. Future research could expand the participant pool to international contexts and test the CHAT-RV model longitudinally to assess the sustainability of its effects. In this way, the study provides a vital starting point for strengthening digital epistemic literacy in the age of artificial intelligence, while opening avenues for further exploration among global researchers and education practitioners.

REFERENCES

- Adarkwah, M. A. (2025). GenAI-Infused Adult Learning in the Digital Era: A Conceptual Framework for Higher Education. *Adult Learning*, 36(3), 149-161.
<https://doi.org/10.1177/10451595241271161>
- Almulla, M. A. (2020). The Effectiveness of the Project-Based Learning (PBL) Approach as a Way to Engage Students in Learning. *SAGE Open*, 10(3), 2158244020938702.
<https://doi.org/10.1177/2158244020938702>

- Biswas, S. (2023). ChatGPT and the Future of Medical Writing. *Radiology*, 307(2), e223312. <https://doi.org/10.1148/radiol.223312>
- Bridges, L. M., McElroy, K., & Welhouse, Z. (2024). Generative Artificial Intelligence: 8 Critical Questions for Libraries. *Journal of Library Administration*, 64(1), 66-79. <https://doi.org/10.1080/01930826.2024.2292484>
- Cacicio, S., & Riggs, R. (2023). ChatGPT: Leveraging AI to Support Personalized Teaching and Learning. *Adult Literacy Education: The International Journal of Literacy, Language, and Numeracy*, 5(2), 70-74. <https://doi.org/10.35847/SCacicio.RRiggs.5.2.70>
- Chiu, T. K. F. (2024). The impact of Generative AI (GenAI) on practices, policies and research direction in education: A case of ChatGPT and Midjourney. *Interactive Learning Environments*, 32(10), 6187-6203. <https://doi.org/10.1080/10494820.2023.2253861>
- Ciampa, K., Wolfe, Z. M., & Bronstein, B. (2023a). ChatGPT in education: Transforming digital literacy practices. *Journal of Adolescent & Adult Literacy*, 67(3), 186-195. <https://doi.org/10.1002/jaal.1310>
- Ciampa, K., Wolfe, Z. M., & Bronstein, B. (2023b). ChatGPT in education: Transforming digital literacy practices. *Journal of Adolescent & Adult Literacy*, 67(3), 186-195. <https://doi.org/10.1002/jaal.1310>
- Creswell, J. W., & Creswell, J. D. (2022). *Research Design*. SAGE Publications, Inc. <https://us.sagepub.com/en-us/nam/research-design/book270550>
- Cronbach, L. J. (1951). Coefficient alpha and the internal structure of tests. *Psychometrika*, 16(3), 297-334. <https://doi.org/10.1007/BF02310555>
- Davison, R. M., Chughtai, H., Nielsen, P., Marabelli, M., Iannacci, F., van Offenbeek, M., Tarafdar, M., Trenz, M., Techatassanasontorn, A. A., Díaz Andrade, A., & Panteli, N. (2024). The ethics of using generative AI for qualitative data analysis. *Information Systems Journal*, 34(5), 1433-1439. <https://doi.org/10.1111/isj.12504>
- Dekov, I. (2025, June 25). The Psychology of Influence: How Social Media Algorithms Fuel Manipulation and Echo Chambers. *VISION_FACTORY*. <https://www.visionfactory.org/post/the-psychology-of-influence-how-social-media-algorithms-fuel-manipulation-and-echo-chambers>
- Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., Baabdullah, A. M., Koohang, A., Raghavan, V., Ahuja, M., Albanna, H., Albashrawi, M. A., Al-Busaidi, A. S., Balakrishnan, J., Barlette, Y., Basu, S., Bose, I., Brooks, L., Buhalis, D., ... Wright, R. (2023). "So what if ChatGPT wrote it?" Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *International Journal of Information Management*, 71, Scopus. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- Ecker, P. A. (2025). Algorithms, Polarization, and the Digital Age: A Literature Review. In P. A. Ecker (Ed.), *The Digital Reinforcement of Bias and Belief: Understanding the Cognitive and Social Impact of Web-Based Information Processing* (pp. 21-43). Springer Nature Switzerland. https://doi.org/10.1007/978-3-031-89998-0_2
- Floridi, L. (2023). *The Ethics of Artificial Intelligence: Principles, Challenges, and Opportunities*. Oxford University Press. <https://doi.org/10.1093/oso/9780198883098.001.0001>
- González-Pérez, L. I., Montoya, M. S. R., & García-Peñalvo, F. J. (2022). *Habilitadores tecnológicos 4.0 para impulsar la educación abierta: Aportaciones para las recomendaciones de la UNESCO. RIED-Revista Iberoamericana de Educación a Distancia*, 25(2), 23-48. <https://doi.org/10.5944/ried.25.2.33088>
- Hamed, A. A., Zachara-Szymanska, M., & Wu, X. (2024). Safeguarding authenticity for mitigating the harms of generative AI: Issues, research agenda, and policies for detection, fact-checking, and ethical AI. *iScience*, 27(2). <https://doi.org/10.1016/j.isci.2024.108782>
- Hayashi, K., & Yuan, K. H. (2023). On the Relationship Between Coefficient Alpha and Closeness Between Factors and Principal Components for the Multi-factor Model. *Springer Proc. Math. Stat.*, 422, 173-185. https://doi.org/10.1007/978-3-031-27781-8_16
- Knowles, M. S., III, E. F. H., & Swanson, R. A. (2014). *The Adult Learner: The definitive classic in adult education and human resource development* (8th ed.). Routledge.
- Kreps, S., & Kriner, D. (2023). How AI Threatens Democracy. *Journal of Democracy*, 34(4), 122-131. <https://doi.org/10.1353/jod.2023.a907693>
- Kurt, D. S. (2020, January 8). Problem-Based Learning (PBL). *Educational Technology*. <https://educationaltechnology.net/problem-based-learning-pbl/>
- Lund, B. D., & Wang, T. (2023). Chatting about ChatGPT: How may AI and GPT impact academia and libraries? *Library Hi Tech News*, 40(3), 26-29. <https://doi.org/10.1108/LHTN-01-2023-0009>
- Mishra, P., Warr, M., & Islam, R. (2023). TPACK in the age of ChatGPT and Generative AI. *Journal of Digital Learning in Teacher Education*, 39(4), 235-251. <https://doi.org/10.1080/21532974.2023.2247480>
- Nurhayati, S., Taufikin, T., Judijanto, L., & Musa, S. (2025). Towards Effective Artificial Intelligence-Driven Learning in Indonesian Child Education: Understanding Parental Readiness, Challenges, and Policy Implications. *Educational Process: International Journal*. <https://www.edupij.com/index/arsiv/76/526/towards-effective-artificial-intelligence-driven-learning-in-indonesian-child-education-understanding-parental-readiness-challenges-and-policy-implications> <https://doi.org/10.22521/edupij.2025.15.155>
- Salaverría, R., & Cardoso, G. (2023). Future of disinformation studies: Emerging research fields. *Profesional de La Información*, 32(5), Article 5. <https://doi.org/10.3145/epi.2023.sep.25>
- Shah, S. B., Thapa, S., Acharya, A., Rauniyar, K., Poudel, S., Jain, S., Masood, A., & Naseem, U. (2024a). Navigating the Web of Disinformation and Misinformation: Large Language Models as Double-Edged Swords. *IEEE Access*, 1-1. <https://doi.org/10.1109/ACCESS.2024.3406644>
- Shah, S. B., Thapa, S., Acharya, A., Rauniyar, K., Poudel, S., Jain, S., Masood, A., & Naseem, U. (2024b). Navigating the Web of Disinformation and Misinformation: Large Language Models as Double-Edged Swords. *IEEE Access*, 1-1. <https://doi.org/10.1109/ACCESS.2024.3406644>
- Shu, K., Sliva, A., Wang, S., Tang, J., & Liu, H. (2017). Fake News Detection on Social Media: A Data Mining Perspective (No. arXiv:1708.01967). arXiv. <https://doi.org/10.1145/3137597.3137600>
- Stark, L. (2023). Breaking Up (with) AI Ethics. *American Literature*, 95(2), 365-379. <https://doi.org/10.1215/00029831-10575148>
- Taufikin, T., Judijanto, L., & Nurhayati, S. (2025). Enhancing Undergraduate Research Writing Using ChatGPT: Effectiveness, Student Perceptions, and Ethical Implications. *Jurnal Edutech Undiksha*, 13(1), 148-157. <https://doi.org/10.23887/jeu.v13i1.94305>
- Thorp, H. H. (2024). ChatGPT to the rescue? *Science*, 385(6714), 1143-1143. <https://doi.org/10.1126/science.adt0007>
- Tlili, A., Shehata, B., Adarkwah, M. A., Bozkurt, A., Hickey, D. T., Huang, R., & Agyemang, B. (2023). What if the devil is my guardian angel: ChatGPT as a case study of using chatbots in education. *Smart Learning Environments*, 10(1), 15. <https://doi.org/10.1186/s40561-023-00237-x>
- UNESCO. (2023). *Guidance for generative AI in education and research*. United Nations Educational, Scientific and Cultural Organization. <https://unesdoc.unesco.org/ark:/48223/pf0000386693>
- Wineburg, S., & McGrew, S. (2019). *Lateral Reading and the Nature of Expertise: Reading Less and Learning More When Evaluating Digital Information*. Teachers College Record,

121(11), 1-40.

<https://doi.org/10.1177/016146811912101102>

Zhou, S., & Yang, S. (2024). The Advantages of Introducing Multimedia Technology into College Students' Physical Education. Lecture Notes of the Institute for Computer

Sciences, Social-Informatics and Telecommunications Engineering, LNICST, 584 LNICST, 61-71. Scopus.

https://doi.org/10.1007/978-3-031-63142-9_6

© 2025 Taufikin

This is an open-access article licensed under the terms of the Creative Commons Attribution License (<http://creativecommons.org/licenses/by/4.0/>), which permits unrestricted use, distribution, and reproduction in any medium, provided the work is properly cited.