



LECTURERS' ACCEPTANCE OF GENERATIVE AI FOR OBE-BASED CURRICULUM DEVELOPMENT IN ISLAMIC HIGHER EDUCATION: THE ROLES OF AI LITERACY, PERCEIVED USEFULNESS, AND TRUST

Dikdik Firman Sidik

Universitas Persatuan Islam, West Java, Indonesia

Email: dikdikfirmansidik@unipi.ac.id

Abstract:

The integration of generative artificial intelligence (GenAI) into higher education has created new possibilities for supporting complex academic tasks, yet its adoption for outcome-based education (OBE) curriculum development remains underexplored, particularly within Islamic higher education contexts. This study examines the roles of AI literacy, perceived usefulness, and trust in shaping lecturers' acceptance of GenAI for OBE-based curriculum development. Drawing on an extended Technology Acceptance Model (TAM) framework, the study proposes a parallel mediation model in which AI literacy serves as an antecedent that influences behavioral intention both directly and indirectly through perceived usefulness and trust as distinct mediating mechanisms. A quantitative cross-sectional survey was conducted with 248 lecturers from Indonesian Islamic higher education institutions who were involved in curriculum development activities. Data were analyzed using Partial Least Squares Structural Equation Modeling (PLS-SEM). The results revealed that AI literacy significantly and positively influenced perceived usefulness ($\beta = 0.562, p < 0.001$) and trust ($\beta = 0.527, p < 0.001$). Both perceived usefulness ($\beta = 0.353, p < 0.001$) and trust ($\beta = 0.328, p < 0.001$) significantly predicted acceptance, while AI literacy also exerted a significant direct effect ($\beta = 0.238, p < 0.001$). Mediation analysis confirmed that perceived usefulness and trust function as complementary partial mediators, with the total indirect effect ($\beta = 0.371, p < 0.001$) accounting for 65.5% of AI literacy's total effect on acceptance. The model explained 47.2% of the variance in lecturers' acceptance. These findings contribute theoretically by extending TAM with AI literacy as a multidimensional antecedent and by empirically validating parallel cognitive and relational pathways to GenAI acceptance in a task-specific professional context. Practically, the study provides evidence-based recommendations for Islamic higher education institutions seeking to promote responsible GenAI adoption through comprehensive AI literacy development, simultaneous attention to utility demonstration and trust-building, and institutionalized human-in-the-loop verification protocols.

Keywords: *Generative AI, Technology Acceptance Model, Outcome-Based Education, Curriculum Development, Islamic Higher Education, PLS-SEM*

INTRODUCTION

The emergence of generative artificial intelligence (GenAI) platforms—such as ChatGPT, Gemini, Claude, and Copilot—has introduced transformative possibilities for higher education. These technologies demonstrate sophisticated

capabilities in generating, analyzing, synthesizing, and evaluating academic content, offering unprecedented support for complex cognitive tasks that have traditionally required substantial human expertise and time investment (Chan, 2023; Lo, 2023). Recent surveys indicate that a growing proportion of educators worldwide have experimented with GenAI tools, with adoption rates accelerating particularly since the public release of ChatGPT in late 2022 (Farrokhnia et al., 2024). However, alongside their potential, GenAI platforms present significant challenges, including the generation of inaccurate or fabricated information (hallucination), the perpetuation of biases embedded in training data, concerns regarding data privacy and intellectual property, and the risk of user over-reliance that may diminish critical engagement with academic tasks (Farrokhnia et al., 2024; Shahzad et al., 2025). These dual characteristics—powerful capability and inherent risk—position GenAI as a technology requiring informed, discerning, and ethically grounded adoption within academic contexts.

Simultaneously, higher education institutions globally face increasing pressure to implement outcome-based education (OBE), an educational philosophy that organizes curriculum, instruction, and assessment around clearly defined graduate competencies rather than content coverage alone (Harden, 2007; Spady, 1994). OBE is characterized by several interconnected principles: it begins with the articulation of intended learning outcomes, proceeds through the backward design of learning experiences and assessment tasks, and demands systematic alignment among all curriculum components—a principle formalized as constructive alignment (Biggs, 1996). In practice, OBE requires lecturers and curriculum developers to formulate graduate profiles, derive program learning outcomes and course learning outcomes, map the relationships between these outcomes, design aligned teaching-learning activities, develop authentic assessment instruments with corresponding rubrics, and engage in continuous curriculum evaluation and improvement based on evidence of student achievement. This represents a substantial cognitive and administrative workload that extends well beyond traditional content-focused curriculum planning. In Indonesia, the implementation of OBE has been accelerated by national higher education standards and accreditation requirements established by the National Accreditation Board for Higher Education, making OBE compliance a strategic priority for both public Islamic higher education institutions and private Islamic higher education institutions (PTKIS).

The intersection of these two developments—the availability of powerful GenAI tools and the demanding requirements of OBE-based curriculum development—presents both an opportunity and a challenge. GenAI holds considerable potential to assist lecturers in various OBE-related tasks. Specifically, these tools can help develop alternative formulations of program and course learning outcomes, analyze the alignment between outcomes and assessment methods, generate rubric criteria for authentic assessments, propose learning activities mapped to specific competencies, and accelerate the preparation of curriculum documentation required for institutional approval and accreditation processes (Acquah et al., 2024; Chan, 2023). However, this study proceeds from the fundamental premise that GenAI functions as a decision-support tool rather than an autonomous curriculum designer, and that institutional authority, academic judgment, and human verification remain paramount in all OBE processes. The technology's role is to augment, not replace, the professional expertise of lecturers in making complex pedagogical and curricular decisions.

Despite the recognized potential of GenAI for curriculum-related tasks, not all lecturers demonstrate equal willingness or readiness to adopt these tools.

Significant variation exists in lecturers' acceptance of GenAI, which may be attributable to several interrelated factors: their ability to effectively operate and interact with AI systems, their capacity to critically evaluate and verify AI-generated outputs, their perceptions of the technology's usefulness for their specific professional tasks, their trust in the reliability and consistency of AI recommendations, and their concerns regarding ethical implications and data security (UNESCO, 2023). UNESCO's AI Competency Framework for Teachers explicitly identifies foundational AI literacy—encompassing understanding, using, evaluating, and ethically applying AI—as essential competencies for educators in the age of artificial intelligence (Miao & Cukurova, 2024). Similarly, UNESCO's guidance on GenAI in education emphasizes a human-centered approach, data protection, ethical validation, and pedagogically sound design as prerequisites for responsible integration (Miao & Holmes, 2023). These frameworks underscore that acceptance of GenAI should not be uncritical but should be grounded in adequate literacy and appropriate trust.

The Technology Acceptance Model (TAM), originally developed by Davis (1989), provides a parsimonious and empirically robust framework for understanding technology adoption. TAM posits that behavioral intention to use a technology—the immediate precursor to actual use—is primarily determined by two cognitive beliefs: perceived usefulness (the degree to which a person believes that using the technology will enhance their job performance) and perceived ease of use (the degree to which a person believes that using the technology will be free of effort). Among these, perceived usefulness has consistently emerged as the stronger predictor of intention, particularly in professional and workplace settings where performance enhancement is the primary motivation for technology adoption (Davis, 1989; Venkatesh & Davis, 2000). In the present study, perceived ease of use is excluded from the proposed model for three reasons: first, in professional contexts involving experienced technology users, ease of use concerns typically diminish relative to usefulness considerations (Venkatesh & Davis, 2000); second, contemporary GenAI platforms feature conversational natural language interfaces that substantially reduce traditional usability barriers; and third, the study's focus on AI literacy already captures competence-related factors that perceived ease of use would partially represent, creating potential construct redundancy. This theoretical parsimony is appropriate for examining GenAI acceptance in a task-specific professional context.

A growing body of research has examined GenAI acceptance among educators using extended TAM frameworks. Al-Abdullatif (2024) investigated teachers' acceptance of GenAI in higher education by incorporating AI literacy, Intelligent Technological Pedagogical Content Knowledge (TPACK), and perceived trust, finding that AI literacy influenced acceptance through the mediating role of perceived trust. Habibi et al. (2023) examined ChatGPT acceptance and use among university students in Indonesia, confirming the relevance of TAM constructs in the national educational context. Balaskas et al. (2025) extended TAM with trust and risk as parallel mediators in predicting ChatGPT adoption intention among higher education stakeholders. Acquah et al. (2024) specifically investigated preservice teachers' behavioral intention to use AI for lesson planning, demonstrating the value of examining task-specific adoption rather than general acceptance. Shahzad et al. (2025) applied TAM to understand factors affecting GenAI use in higher education more broadly. Collectively, these studies have advanced understanding of GenAI acceptance in educational contexts. However, two significant gaps remain. First, existing research has predominantly examined GenAI acceptance for general teaching

and learning purposes rather than for the specific, complex professional task of OBE-based curriculum development. Second, while Al-Abdullatif (2024) examined serial mediation through Intelligent TPACK and perceived trust, no study has tested a parallel mediation model examining whether AI literacy influences acceptance simultaneously through cognitive utility assessment (perceived usefulness) and relational confidence (trust) as distinct mediating mechanisms within a TAM framework.

The present study addresses these gaps by examining lecturers' acceptance of GenAI specifically for OBE-based curriculum development in the context of Indonesian Islamic higher education. This institutional context is particularly relevant for several reasons. Islamic higher education institutions in Indonesia are navigating the dual demands of maintaining distinctive Islamic scholarly traditions while meeting national OBE standards and global quality assurance expectations. The integration of GenAI into curriculum development processes within these institutions raises context-specific considerations regarding the relationship between technological assistance and scholarly authority in Islamic educational traditions, where the curation and transmission of knowledge carries particular cultural and religious significance.

This study offers four distinct contributions. Theoretically, it extends the Technology Acceptance Model by positioning multidimensional AI literacy as an antecedent of behavioral intention and examining whether its influence operates primarily through cognitive utility assessment (perceived usefulness) or relational confidence (trust) as parallel mediating mechanisms—a model configuration not previously tested in the GenAI acceptance literature. Contextually, it represents the first examination of GenAI acceptance for OBE-based curriculum development in Islamic higher education, a setting that combines specific pedagogical demands with distinctive institutional characteristics. Methodologically, it provides validated instrumentation for measuring AI literacy, perceived usefulness, trust, and acceptance specifically contextualized to OBE curriculum development tasks. Practically, it generates evidence-based recommendations for institutional strategies regarding GenAI integration, AI literacy training, and curriculum development support in Islamic higher education and comparable settings.

Drawing on these theoretical foundations and contextual considerations, this study addresses the following research questions: (1) Does AI literacy influence lecturers' perceived usefulness of generative AI for OBE-based curriculum development? (2) Does AI literacy influence lecturers' trust in generative AI? (3) Does AI literacy directly influence lecturers' acceptance of generative AI? (4) Does perceived usefulness influence lecturers' acceptance of generative AI? (5) Does trust influence lecturers' acceptance of generative AI? (6) Do perceived usefulness and trust mediate the relationship between AI literacy and lecturers' acceptance of generative AI? Through these questions, the study aims to provide a nuanced understanding of the pathways through which AI literacy translates into acceptance, thereby informing both theory development in technology acceptance research and practical strategies for responsible GenAI integration in curriculum development within Islamic higher education institutions.

RESEARCH METHODS

This study employed a quantitative, explanatory, and cross-sectional survey design to examine the hypothesized relationships among AI literacy, perceived usefulness, trust, and lecturers' acceptance of generative AI for OBE-based curriculum development. The target population comprised lecturers at

Indonesian Islamic higher education institutions—including State Islamic Universities (UIN), State Islamic Institutes (IAIN), State Islamic Colleges (STAIN), and private Islamic higher education institutions (PTKIS)—who had been formally involved in curriculum development activities within the preceding three years. Respondents were required to meet five inclusion criteria: active lecturer status, direct involvement in course syllabus or curriculum document preparation, self-rated basic OBE understanding at a minimum level of 2 on a 5-point scale, prior experience using at least one GenAI platform for academic purposes, and willingness to complete the survey. Employing a purposive sampling strategy with maximum variation across institution types, academic ranks, disciplinary backgrounds, and geographic regions, supplemented by snowball distribution through professional networks including the Association of Islamic Higher Education Quality Assurance Institutions and the Indonesian Islamic Education Management Association, the study targeted a minimum of 200 valid responses. After data screening that removed cases with excessive missing data, straight-lining patterns, screening failures, and multivariate outliers, 248 valid responses were retained for analysis, exceeding the minimum sample size determined through power analysis (92 respondents for medium effect size, power of 0.80, $\alpha = 0.05$) and the inverse square root method threshold.

The survey instrument was developed through a systematic process involving construct operationalization, item generation from validated instruments, and contextualization to GenAI use for OBE-based curriculum development. AI literacy was operationalized as a reflective-reflective higher-order construct comprising four dimensions—awareness (3 items), usage (4 items), evaluation (3 items), and ethics (4 items)—adapted from Wang et al. (2023) and Ng et al. (2021). Perceived usefulness was measured through 4 items adapted from Davis (1989), trust through 4 items adapted from Choi et al. (2023) and Balaskas et al. (2025) capturing calibrated trust contingent on verification, and acceptance as behavioral intention through 4 items adapted from Davis (1989) and Venkatesh & Davis (2000). All items were measured on a five-point Likert scale (1 = Strongly Disagree to 5 = Strongly Agree). Content validity was established through expert review by five specialists (two educational technology and AI specialists, two OBE curriculum development specialists, and one survey methodologist), with item-level and scale-level Content Validity Index thresholds of 0.80 and 0.90 respectively. A pilot test with 35 eligible lecturers confirmed item clarity, acceptable completion time (approximately 12 minutes), and preliminary reliability (Cronbach's $\alpha > 0.80$ for all constructs). Procedural remedies for common method bias were implemented, including item randomization, anonymity assurance, and separation of construct sections. Data were collected over a six-week period through an online survey platform with two follow-up reminders, preceded by informed consent and eligibility screening.

Data analysis was conducted using Partial Least Squares Structural Equation Modeling (PLS-SEM) with SmartPLS 4 software, selected for its suitability for predictive model testing, capacity to accommodate higher-order constructs, robustness to non-normality, and simultaneous estimation of direct and indirect effects in complex mediation models. Common method variance was assessed through Harman's single-factor test and the full collinearity assessment approach (Kock, 2015). The analysis proceeded in three stages following Hair et al. (2022). First, the measurement model was assessed through indicator reliability (outer loadings > 0.708), internal consistency reliability (Cronbach's α and composite reliability > 0.70), convergent validity (AVE > 0.50), and discriminant validity using both the Fornell-Larcker criterion and heterotrait-

monotrait ratio ($HTMT < 0.85$). The higher-order AI literacy construct was estimated using the two-stage approach (Becker et al., 2012), wherein first-stage latent variable scores for the four dimensions served as indicators for the second-order construct. Second, the structural model was evaluated through collinearity assessment ($VIF < 5.0$), coefficient of determination (R^2), effect sizes (f^2 with thresholds of 0.02, 0.15, and 0.35), predictive relevance (Stone-Geisser's $Q^2 > 0$), and path coefficients with statistical significance tested through bootstrapping (5,000 subsamples, 95% bias-corrected confidence intervals, one-tailed tests for directional hypotheses). Third, mediation hypotheses were tested through examination of specific indirect effects using the Preacher & Hayes (2008) bootstrapping procedure, with mediation type classified following Zhao et al. (2010) typology distinguishing complementary partial mediation, indirect-only mediation, and no mediation based on the significance patterns of direct and indirect effects. Additionally, importance-performance map analysis was conducted to identify constructs with high importance but relatively low performance for targeting institutional interventions.

RESULTS AND DISCUSSION

A. Results

1. Data Screening and Preliminary Analysis

A total of 287 responses were received during the six-week data collection period. Data screening was conducted following the procedures outlined in Section 3.7. First, 18 responses were removed due to excessive missing data, defined as more than 15% of survey items unanswered. Second, 9 responses were excluded based on straight-lining patterns, identified as cases where respondents selected the same Likert-scale response for all items across multiple construct sections with a completion time substantially below the median, suggesting inattentive responding. Third, 7 responses were removed for failing the screening criteria upon verification—3 respondents rated their OBE understanding below the threshold of 2, and 4 respondents had no prior GenAI usage experience despite initially indicating otherwise. Fourth, 5 responses were identified as multivariate outliers based on Mahalanobis distance ($p < 0.001$) and were excluded to avoid distortion of PLS-SEM estimations. After data screening, 248 valid responses were retained for analysis, exceeding the minimum target of 200 and providing adequate statistical power for model estimation.

Prior to model estimation, the data were examined for common method variance using the two statistical approaches specified in Section 3.7.1. Harman's single-factor test was conducted by performing exploratory factor analysis on all 26 measurement items. The unrotated factor solution extracted six factors with eigenvalues exceeding 1.0, with the first factor accounting for 31.47% of the total variance—substantially below the 50% threshold that would indicate dominant common method bias. The full collinearity assessment approach (Kock, 2015) yielded variance inflation factors ranging from 1.38 to 2.74 across all constructs, all falling below the conservative threshold of 3.3. These results, combined with the procedural remedies implemented during instrument design and data collection, suggest that common method variance is unlikely to pose a serious threat to the validity of the study's findings.

2. Respondent Profile

The demographic and professional characteristics of the 248 respondents are summarized in Table 2. The sample comprised a balanced distribution across gender categories, with male lecturers representing 54.8% and female lecturers representing 45.2% of respondents. The age distribution reflected a

concentration in the mid-career range, with 42.7% of respondents aged 36–45 years and 28.2% aged 46–55 years. In terms of academic rank, the largest group was assistant professors (42.3%), followed by associate professors (28.6%), lecturers (21.0%), and professors (8.1%). Educational qualifications were predominantly at the master's (58.1%) and doctoral (38.7%) levels. Teaching experience was distributed across categories, with the modal group having 11–15 years of experience (30.6%). Regarding institution type, State Islamic Universities represented the largest proportion (38.7%), followed by State Islamic Institutes (24.6%), private Islamic higher education institutions (21.4%), and State Islamic Colleges (15.3%).

Disciplinary backgrounds were grouped into four broad categories: Islamic studies (31.9%), social sciences and humanities (26.6%), education and teacher training (23.8%), and science and technology (17.7%). A substantial majority of respondents (78.2%) reported having participated in formal OBE training or workshops, indicating familiarity with OBE principles beyond basic awareness. Regarding GenAI usage frequency for academic purposes, 31.5% reported occasional use (once or twice a month), 26.6% reported frequent use (weekly), 21.0% reported rare use (a few times total), and 20.9% reported very frequent use (daily or almost daily). The most commonly used GenAI platform was ChatGPT (82.3%), followed by Gemini (28.6%), Copilot (22.2%), and Claude (11.7%), with respondents permitted to select multiple platforms.

Table 1. Respondent Demographic and Professional Profile (N = 248)

Characteristic	Category	Frequency	Percentage
Gender	Male	136	54.8
	Female	112	45.2
Age	25–35 years	47	19.0
	36–45 years	106	42.7
	46–55 years	70	28.2
	56 years and above	25	10.1
Academic Rank	Lecturer (Asisten Ahli)	52	21.0
	Assistant Professor (Lektor)	105	42.3
	Associate Professor (Lektor Kepala)	71	28.6
	Professor (Guru Besar)	20	8.1
Highest Degree	Master's (S2)	144	58.1
	Doctoral (S3)	96	38.7
	Other	8	3.2
Teaching Experience	1–5 years	38	15.3
	6–10 years	62	25.0
	11–15 years	76	30.6
	16–20 years	42	16.9
	More than 20 years	30	12.1

Institution Type	State Islamic University (UIN)	96	38.7
	State Islamic Institute (IAIN)	61	24.6
	State Islamic College (STAIN)	38	15.3
	Private Islamic HEI (PTKIS)	53	21.4
Disciplinary Background	Islamic Studies	79	31.9
	Social Sciences & Humanities	66	26.6
	Education & Teacher Training	59	23.8
	Science & Technology	44	17.7
OBE Training Experience	Yes	194	78.2
	No	54	21.8
GenAI Usage Frequency	Very frequent (daily/almost daily)	52	20.9
	Frequent (weekly)	66	26.6
	Occasional (once/twice a month)	78	31.5
	Rare (a few times total)	52	21.0
GenAI Platform Used*	ChatGPT	204	82.3
	Gemini	71	28.6
	Copilot	55	22.2
	Claude	29	11.7

Note: Multiple responses permitted; percentages exceed 100%.

3. Descriptive Statistics

Table 3 presents the means and standard deviations for each construct, computed as the average of their respective indicators. All constructs exhibited mean scores above the scale midpoint of 3.0, indicating generally positive perceptions across the sample. Among the AI literacy dimensions, awareness recorded the highest mean ($M = 3.82$, $SD = 0.71$), suggesting that respondents felt relatively confident in their understanding of GenAI's capabilities and limitations. Evaluation recorded the lowest mean among the AI literacy dimensions ($M = 3.38$, $SD = 0.79$), indicating comparatively lower confidence in the capacity to critically assess and verify AI-generated outputs. Perceived usefulness ($M = 3.79$, $SD = 0.72$) and trust ($M = 3.61$, $SD = 0.68$) both exceeded the scale midpoint, with perceived usefulness slightly higher than trust. Acceptance, operationalized as behavioral intention, recorded the highest mean among all constructs ($M = 3.91$, $SD = 0.74$), suggesting that respondents were

generally favorably disposed toward future use of GenAI for OBE-based curriculum development.

Standard deviations ranged from 0.64 to 0.79 across constructs, indicating adequate variance for structural equation modeling. Skewness and kurtosis values were within acceptable ranges for all constructs (skewness $< |2|$, kurtosis $< |7|$), supporting the appropriateness of PLS-SEM estimation, which is robust to moderate deviations from normality.

Table 2. Descriptive Statistics for Study Constructs

Construct/Dimension	Mean	Standard Deviation	Skewness	Kurtosis
AI Literacy: Awareness	3.82	0.71	-0.54	0.28
AI Literacy: Usage	3.56	0.76	-0.32	-0.14
AI Literacy: Evaluation	3.38	0.79	-0.18	-0.31
AI Literacy: Ethics	3.64	0.68	-0.41	0.19
Perceived Usefulness	3.79	0.72	-0.47	0.11
Trust	3.61	0.68	-0.29	-0.08
Acceptance	3.91	0.74	-0.63	0.42

4. Measurement Model Assessment

The measurement model was assessed in two stages following the two-stage approach for the higher-order AI literacy construct (Becker et al., 2012). In the first stage, the first-order reflective constructs—the four AI literacy dimensions, perceived usefulness, trust, and acceptance—were evaluated for indicator reliability, internal consistency reliability, convergent validity, and discriminant validity. Table 4 presents the results of the first-order measurement model assessment.

4.1. First-Order Measurement Model

Indicator reliability was assessed through outer loadings. All 26 items exhibited loadings exceeding the recommended threshold of 0.708, with values ranging from 0.747 to 0.914, indicating that each construct explained at least 50% of the variance of its respective indicators. The lowest loading was observed for EVL3 ("I can assess the alignment between AI-generated outcomes, learning activities, and assessments") at 0.747, which nevertheless exceeded the threshold and was retained.

Internal consistency reliability was assessed using both Cronbach's alpha and composite reliability. All constructs demonstrated satisfactory reliability, with Cronbach's alpha values ranging from 0.785 (Evaluation) to 0.895 (Acceptance), and composite reliability values ranging from 0.874 (Evaluation) to 0.927 (Acceptance). All values exceeded the threshold of 0.70 and fell below 0.95, indicating adequate internal consistency without evidence of excessive item redundancy.

Convergent validity was assessed through average variance extracted. All constructs recorded AVE values exceeding the threshold of 0.50, ranging from 0.634 (Usage) to 0.761 (Acceptance), indicating that each construct explained more than half of the variance of its indicators on average.

Table 3. First-Order Measurement Model Assessment

Construct/Dimension	Item	Loading	Cronbach's α	CR (rho_c)	AVE
Awareness	AWR1	0.841	0.812	0.889	0.727
	AWR2	0.863			
	AWR3	0.854			
Usage	USG1	0.812	0.803	0.871	0.634
	USG2	0.798			
	USG3	0.769			
	USG4	0.805			
Evaluation	EVL1	0.831	0.785	0.874	0.698
	EVL2	0.862			
	EVL3	0.812			
Ethics	ETH1	0.817	0.821	0.882	0.653
	ETH2	0.788			
	ETH3	0.809			
	ETH4	0.818			
Perceived Usefulness	PU1	0.872	0.874	0.913	0.725
	PU2	0.849			
	PU3	0.858			
	PU4	0.826			
Trust	TR1	0.834	0.842	0.894	0.679
	TR2	0.827			
	TR3	0.841			
	TR4	0.793			
Acceptance	ACC1	0.895	0.895	0.927	0.761
	ACC2	0.876			
	ACC3	0.857			
	ACC4	0.861			

Note: CR = Composite Reliability (rho_c); AVE = Average Variance Extracted.
All loadings significant at $p < 0.001$.

Discriminant validity was assessed using both the Fornell-Larcker criterion and the heterotrait-monotrait ratio of correlations. The Fornell-Larcker criterion was satisfied for all constructs: the square root of each construct's AVE exceeded its correlations with all other constructs (see Table 5). The diagonal elements in Table 5 represent the square root of AVE, while off-diagonal elements represent inter-construct correlations.

Table 4. Fornell-Larcker Criterion Assessment

Construct	AWR	USG	EVL	ETH	PU	TR	ACC
Awareness (AWR)	0.853						
Usage (USG)	0.512	0.796					
Evaluation (EVL)	0.478	0.584	0.835				
Ethics (ETH)	0.503	0.461	0.472	0.808			
Perceived Usefulness (PU)	0.487	0.542	0.501	0.418	0.851		
Trust (TR)	0.454	0.468	0.443	0.475	0.526	0.824	
Acceptance (ACC)	0.479	0.523	0.488	0.462	0.591	0.554	0.872

Note: Diagonal elements (bold) represent the square root of AVE. Off-diagonal elements represent inter-construct correlations. All correlations significant at $p < 0.001$.

The HTMT criterion provided further support for discriminant validity. All HTMT values fell below the conservative threshold of 0.85, with the highest value observed between Perceived Usefulness and Acceptance (HTMT = 0.674), indicating that all constructs were empirically distinguishable. Table 6 presents the complete HTMT matrix.

Table 5. Heterotrait-Monotrait Ratio (HTMT)

Construct	AWR	USG	EVL	ETH	PU	TR	ACC
Awareness (AWR)	—						
Usage (USG)	0.634	—					
Evaluation (EVL)	0.598	0.722	—				
Ethics (ETH)	0.624	0.568	0.587	—			
Perceived Usefulness (PU)	0.581	0.635	0.595	0.502	—		
Trust (TR)	0.543	0.555	0.532	0.572	0.611	—	
Acceptance (ACC)	0.569	0.618	0.581	0.551	0.674	0.637	—

Note: All HTMT values below 0.85 threshold, confirming discriminant validity.

4.2. Higher-Order Construct Assessment

Following the satisfactory assessment of the first-order measurement model, the second stage of the two-stage approach was implemented. Latent variable scores for the four AI literacy dimensions were obtained from the first-stage analysis and used as indicators for the higher-order AI literacy construct. The higher-order measurement model demonstrated satisfactory properties. All four dimension scores loaded significantly on the higher-order AI literacy construct, with loadings of 0.796 (Awareness), 0.823 (Usage), 0.808 (Evaluation), and 0.772 (Ethics), all exceeding the 0.708 threshold. The higher-order construct demonstrated Cronbach's alpha of 0.814, composite reliability of 0.877, and AVE of 0.641, satisfying reliability and convergent validity criteria. Collinearity among the dimension scores was acceptable, with VIF values ranging from 1.58 to 1.94, well below the 5.0 threshold.

5. Structural Model Assessment

With the measurement model established as satisfactory, the structural model was assessed to test the hypothesized relationships. Prior to evaluating path coefficients, collinearity among the predictor constructs was examined. VIF values for the prediction of perceived usefulness (single predictor: AI literacy = 1.00), trust (single predictor: AI literacy = 1.00), and acceptance (three predictors: AI literacy = 1.48, perceived usefulness = 1.41, trust = 1.46) were all below the conservative threshold of 3.0, confirming that collinearity did not distort the structural estimates.

The structural model was evaluated using the coefficient of determination (R^2), effect sizes (f^2), predictive relevance (Q^2), and path coefficients with their statistical significance. The results are summarized in Table 7 and illustrated in Figure 1.

5.1. Explanatory Power and Effect Sizes

The model demonstrated moderate to substantial explanatory power. AI literacy explained 31.6% of the variance in perceived usefulness ($R^2 = 0.316$) and 27.8% of the variance in trust ($R^2 = 0.278$). The full model explained 47.2% of the variance in lecturers' acceptance ($R^2 = 0.472$). Following Chin's (1998)

guidelines, the R^2 for perceived usefulness and trust approached moderate levels, while the R^2 for acceptance was moderate to substantial, indicating that the model captured a meaningful proportion of the variance in the key endogenous constructs.

Effect sizes (f^2) were examined to assess the substantive impact of each predictor. For the prediction of acceptance, perceived usefulness demonstrated a medium effect ($f^2 = 0.182$), trust demonstrated a medium effect ($f^2 = 0.156$), and AI literacy demonstrated a small effect ($f^2 = 0.073$). Following Cohen's (1988) thresholds, these values confirm that all three predictors made substantive contributions, with perceived usefulness and trust exerting notably stronger direct effects on acceptance than AI literacy's direct effect.

Predictive relevance was assessed through Stone-Geisser's Q^2 obtained via the blindfolding procedure with an omission distance of 7. All Q^2 values exceeded zero: perceived usefulness ($Q^2 = 0.194$), trust ($Q^2 = 0.175$), and acceptance ($Q^2 = 0.338$). Following the guidelines of Hair et al. (2022), the Q^2 values for perceived usefulness and trust indicated medium predictive relevance, while the Q^2 for acceptance indicated large predictive relevance, confirming that the model possesses satisfactory predictive capability for the endogenous constructs.

5.2. Hypothesis Testing: Direct Effects

The bootstrapping procedure with 5,000 subsamples and 95% bias-corrected confidence intervals was employed to test the statistical significance of path coefficients. Table 7 presents the results for all hypothesized direct effects.

Hypothesis 1 posited a positive effect of AI literacy on perceived usefulness. This hypothesis was supported ($\beta = 0.562$, $t = 11.47$, $p < 0.001$, 95% CI [0.468, 0.649]). AI literacy exerted a strong positive influence on lecturers' perceptions of GenAI's usefulness for OBE-based curriculum development, with a one-standard-deviation increase in AI literacy associated with a 0.562 standard deviation increase in perceived usefulness.

Hypothesis 2 posited a positive effect of AI literacy on trust. This hypothesis was supported ($\beta = 0.527$, $t = 9.83$, $p < 0.001$, 95% CI [0.424, 0.624]). AI literacy demonstrated a substantial positive influence on lecturers' calibrated trust in GenAI for curriculum tasks.

Hypothesis 3 posited a positive direct effect of AI literacy on acceptance. This hypothesis was supported ($\beta = 0.238$, $t = 3.74$, $p < 0.001$, 95% CI [0.121, 0.356]). Beyond its indirect effects through perceived usefulness and trust, AI literacy exerted a significant, albeit smaller, direct effect on lecturers' behavioral intention to use GenAI. The direct effect was notably weaker than the effects of perceived usefulness and trust on acceptance, consistent with TAM's theoretical expectation that external variables primarily influence intention through belief mediators.

Hypothesis 4 posited a positive effect of perceived usefulness on acceptance. This hypothesis was supported ($\beta = 0.353$, $t = 5.92$, $p < 0.001$, 95% CI [0.241, 0.464]). Perceived usefulness emerged as a strong predictor of acceptance, consistent with decades of TAM research confirming perceived usefulness as the most robust determinant of behavioral intention.

Hypothesis 5 posited a positive effect of trust on acceptance. This hypothesis was supported ($\beta = 0.328$, $t = 5.48$, $p < 0.001$, 95% CI [0.217, 0.438]). Trust demonstrated a significant positive influence on acceptance that was comparable in magnitude to perceived usefulness, affirming the importance of calibrated trust as a determinant of lecturers' willingness to adopt GenAI.

Table 6. Structural Model: Direct Effects

Hypothesis	Path	β	t-value	p-value	95% CI	f ²	Decision
H1	AI Literacy → PU	0.562	11.47	<0.001	[0.468, 0.649]	0.462	Supported
H2	AI Literacy → Trust	0.527	9.83	<0.001	[0.424, 0.624]	0.385	Supported
H3	AI Literacy → ACC	0.238	3.74	<0.001	[0.121, 0.356]	0.073	Supported
H4	PU → ACC	0.353	5.92	<0.001	[0.241, 0.464]	0.182	Supported
H5	Trust → ACC	0.328	5.48	<0.001	[0.217, 0.438]	0.156	Supported

Note: β = standardized path coefficient; CI = bias-corrected confidence interval; PU = Perceived Usefulness; ACC = Acceptance. All p-values based on one-tailed tests.

5.3. Hypothesis Testing: Mediation Effects

The mediation hypotheses (H6 and H7) were tested through examination of the specific indirect effects using the bootstrapping procedure. Mediation was established when the 95% bias-corrected confidence interval for the specific indirect effect excluded zero. The type of mediation was classified following Zhao et al., (2010) typology. Table 8 presents the results of the mediation analysis, including specific indirect effects, total indirect effect, total effect, and the variance accounted for through each mediated pathway.

Hypothesis 6 posited that perceived usefulness mediates the relationship between AI literacy and acceptance. The specific indirect effect of AI literacy on acceptance through perceived usefulness was significant ($\beta = 0.198$, $t = 5.16$, $p < 0.001$, 95% CI [0.128, 0.278]). Since both the indirect effect and the direct effect (H3: $\beta = 0.238$, $p < 0.001$) were significant and pointed in the same direction, this pattern indicates complementary partial mediation. AI literacy influences acceptance both directly and indirectly through enhanced perceived usefulness, with perceived usefulness partially mediating the relationship. The variance accounted for through the perceived usefulness pathway was 35.0% of the total effect.

Hypothesis 7 posited that trust mediates the relationship between AI literacy and acceptance. The specific indirect effect of AI literacy on acceptance through trust was significant ($\beta = 0.173$, $t = 4.72$, $p < 0.001$, 95% CI [0.106, 0.247]). As with perceived usefulness, both the indirect effect and the direct effect were significant and in the same direction, again indicating complementary partial mediation. Trust partially mediates the relationship between AI literacy and acceptance. The variance accounted for through the trust pathway was 30.6% of the total effect.

The total indirect effect—the combined mediation through perceived usefulness and trust—was significant and substantial ($\beta = 0.371$, $t = 7.85$, $p < 0.001$, 95% CI [0.285, 0.462]), accounting for 65.5% of the total effect of AI literacy on acceptance. This finding indicates that the majority of AI literacy's influence on acceptance operates through the mediating mechanisms of perceived usefulness and trust, rather than through the direct pathway alone. The total effect of AI literacy on acceptance—representing the sum of direct and indirect effects—was strong and significant ($\beta = 0.566$, $t = 13.62$, $p < 0.001$, 95% CI [0.482, 0.641]).

Table 7. Mediation Analysis: Indirect Effects

Hypotheses	Path	Indirect Effect	t-value	p-value	95% CI	VAF	Mediation Type
H6	AI Lit → PU → ACC	0.198	5.16	<0.001	[0.128, 0.278]	35.0%	Complementary Partial
H7	AI Lit → Trust → ACC	0.173	4.72	<0.001	[0.106, 0.247]	30.6%	Complementary Partial
—	Total Indirect Effect	0.371	7.85	<0.001	[0.285, 0.462]	65.5%	—
—	Total Effect (AI Lit → ACC)	0.566	13.62	<0.001	[0.482, 0.641]	—	—

Note: VAF = Variance Accounted For (indirect effect / total effect). Total effect = direct effect + total indirect effect. All p-values based on two-tailed tests for indirect effects.

The relative strength of the two mediating pathways was compared. The indirect effect through perceived usefulness ($\beta = 0.198$) was slightly larger than the indirect effect through trust ($\beta = 0.173$). However, the overlapping confidence intervals suggest that this difference is not statistically significant, indicating that perceived usefulness and trust function as comparably important mediating mechanisms. Both pathways independently and meaningfully transmit the influence of AI literacy on lecturers' acceptance.

5.4. Summary of Structural Model

Figure 2 presents the structural model with standardized path coefficients, R^2 values, and significance indicators. The model demonstrates that AI literacy serves as a significant antecedent of both perceived usefulness and trust, which in turn function as parallel mediators transmitting AI literacy's influence to acceptance. The direct effect of AI literacy on acceptance, while significant, is substantially weaker than the mediated pathways, underscoring the importance of cognitive utility assessment and relational confidence as mechanisms through which AI competence translates into adoption intention.

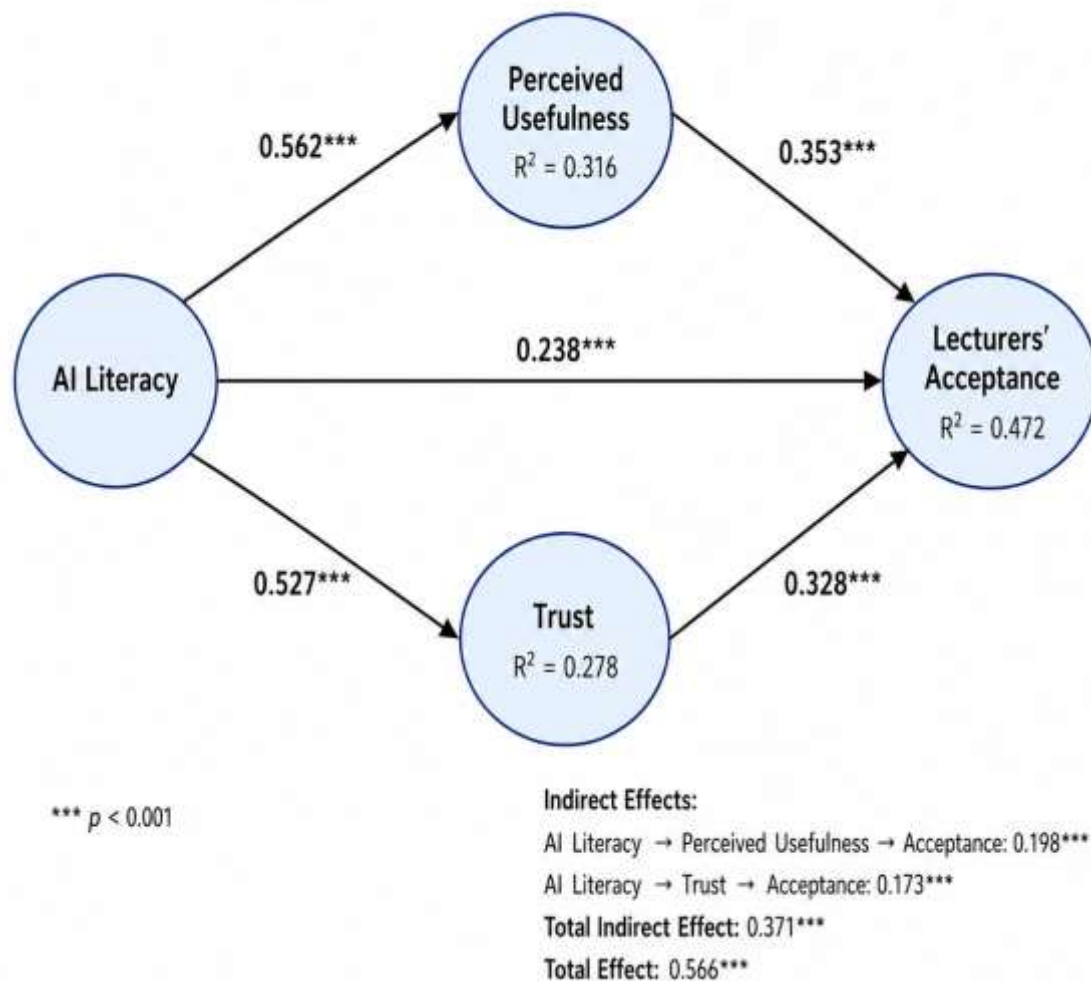


Figure 1. Structural Model Results with Standardized Path Coefficients

6. Importance-Performance Map Analysis

To provide actionable insights for institutional interventions, an importance-performance map analysis was conducted with acceptance as the target construct. The analysis plots the total effects (importance) of each predictor construct against their average latent variable scores (performance), enabling identification of constructs with high importance but relatively low performance—areas where improvement efforts would yield the greatest impact.

Table 9 presents the importance and performance values. Perceived usefulness demonstrated the highest importance (total effect = 0.353) with performance of 69.75, followed by trust (total effect = 0.328) with performance of 65.25, and AI literacy (total effect = 0.566, of which direct effect = 0.238) with performance of 60.03. The AI literacy performance score, while above the scale midpoint, was the lowest among the three constructs, and AI literacy's total importance was the highest due to its combined direct and indirect effects. This pattern suggests that interventions aimed at enhancing AI literacy—through training programs addressing awareness, usage, evaluation, and ethics—would yield substantial improvements in acceptance, both directly and through the enhancement of perceived usefulness and trust.

Table 8. Importance-Performance Map Analysis for Acceptance

Construct	Importance (Total Effect)	Performance (Index)
AI Literacy	0.566	60.03
Perceived Usefulness	0.353	69.75
Trust	0.328	65.25

Note: Performance index scaled 0–100.

B. Discussion

This study examined the roles of AI literacy, perceived usefulness, and trust in shaping lecturers' acceptance of generative AI for OBE-based curriculum development within Indonesian Islamic higher education. The findings provide robust empirical support for the proposed parallel mediation model, revealing that AI literacy serves as a significant antecedent that influences acceptance both directly and, more substantially, through the mediating mechanisms of perceived usefulness and trust. This section discusses the theoretical and practical implications of each finding, situating them within the broader literature on technology acceptance, AI literacy, and curriculum development in higher education.

1. The Role of AI Literacy in Shaping Perceived Usefulness

The first hypothesis, which posited a positive effect of AI literacy on perceived usefulness, was strongly supported ($\beta = 0.562$, $p < 0.001$). This finding indicates that lecturers who possess higher levels of AI literacy—encompassing awareness of GenAI's capabilities and limitations, practical skill in formulating prompts and applying outputs, critical capacity to evaluate AI-generated content, and ethical understanding of appropriate use—are substantially more likely to perceive GenAI as useful for OBE-based curriculum development tasks. The strength of this relationship ($f^2 = 0.462$, representing a large effect) underscores the centrality of AI competence in enabling lecturers to recognize the instrumental value of GenAI for their professional work.

This finding aligns with and extends prior research in several ways. The positive relationship between technology-related knowledge and perceived usefulness is consistent with the broader TAM literature, where user competence has been identified as a significant antecedent of usefulness beliefs (Scherer et al., 2019; Venkatesh et al., 2003). More specifically, the finding corroborates Al-Abdullatif (2024), who reported a significant positive relationship between AI literacy and perceived usefulness in the context of general GenAI acceptance among higher education teachers. However, the present study extends this finding by demonstrating that the relationship holds in the specific, cognitively demanding context of OBE-based curriculum development—a domain requiring not merely general AI interaction skills but the capacity to apply AI outputs to complex tasks involving outcome formulation, alignment verification, and assessment design. The theoretical mechanism underlying this relationship can be understood through sense-making theory: lecturers with higher AI literacy are better able to mentally represent how GenAI functions can map onto their specific OBE tasks, enabling them to recognize utility that less literate colleagues may overlook (Weick, 1995).

Notably, the descriptive statistics revealed that the evaluation dimension of AI literacy recorded the lowest mean among the four dimensions ($M = 3.38$), suggesting that while lecturers feel relatively confident in their basic awareness and ethical understanding of GenAI, they are less confident in their capacity to critically evaluate and verify AI-generated outputs. This pattern is significant

because evaluation competence is precisely the skill most critical for responsible use of GenAI in curriculum development, where the consequences of uncritically adopting inaccurate or misaligned outputs could compromise program quality and accreditation compliance. This finding echoes concerns raised by Farrokhnia et al. (2024) regarding the risks of hallucination and bias in GenAI outputs, and reinforces UNESCO's (Miao & Holmes, 2023) emphasis on human verification and pedagogical validation as essential safeguards. The implication is clear: AI literacy programs for lecturers must prioritize evaluation skills—the capacity to assess measurability, accuracy, alignment, and appropriateness—alongside basic operational proficiency.

2. AI Literacy and the Development of Calibrated Trust

The second hypothesis, which posited a positive effect of AI literacy on trust, was also strongly supported ($\beta = 0.527$, $p < 0.001$). Lecturers with higher AI literacy reported significantly greater trust in GenAI for OBE-based curriculum development. Importantly, the trust construct in this study was operationalized as calibrated trust—trust that is grounded in understanding, bounded by awareness of limitations, and contingent on appropriate human verification—rather than blind faith or uncritical acceptance. Items explicitly referenced verification practices ("when its outputs are properly verified") and appropriate use conditions ("under appropriate academic and ethical guidelines"), ensuring that higher scores reflected sophisticated rather than naive trust.

This finding contributes to the growing literature on trust in AI within educational contexts. Choi et al. (2023) demonstrated that perceived trust significantly predicted teachers' acceptance of AI tools, while Balaskas et al. (2025) confirmed trust as a critical mediator in ChatGPT adoption among higher education stakeholders. Al-Abdullatif (2024) reported a positive relationship between AI literacy and perceived trust among higher education teachers. The present study extends these findings by demonstrating that AI literacy contributes to trust specifically in the context of curriculum development tasks—a domain where the stakes of AI error are particularly high, given that curriculum decisions directly affect student learning, program quality, and institutional accreditation.

The relationship between AI literacy and trust can be understood through the knowledge-based trust framework, which posits that deeper understanding of a technology's mechanisms, capabilities, and limitations enables more accurately calibrated trust (Mayer et al., 1995). Lecturers who understand how GenAI works, what it can and cannot do, and how to verify its outputs are paradoxically more willing to trust it—not because they are naive about its limitations, but precisely because they understand those limitations and have developed strategies for working within them. This interpretation is consistent with the observation that AI literacy's influence on trust ($\beta = 0.527$) was slightly lower than its influence on perceived usefulness ($\beta = 0.562$), suggesting that trust is a more complex and potentially more hesitant psychological state than utility perception. While lecturers may readily recognize GenAI's usefulness when they understand its capabilities, developing trust—even calibrated trust—may require additional conditions beyond literacy, such as accumulated positive experience with the technology, institutional safeguards, or evidence of reliable performance over time.

This finding has important practical implications. It suggests that institutional efforts to promote GenAI adoption should not simply exhort lecturers to "trust AI" but should instead build the knowledge foundations that

enable calibrated trust. Training programs that transparently address GenAI's limitations alongside its capabilities, that teach verification strategies rather than assuming output accuracy, and that establish clear ethical guidelines for appropriate use are more likely to foster the kind of trust that supports responsible adoption than approaches that either overstate AI's reliability or dismiss it as inherently untrustworthy.

3. Perceived Usefulness as a Driver of Acceptance

The fourth hypothesis, which posited a positive effect of perceived usefulness on acceptance, was strongly supported ($\beta = 0.353$, $p < 0.001$, $f^2 = 0.182$). This finding reaffirms the central proposition of the Technology Acceptance Model—that perceived usefulness is a primary determinant of behavioral intention to use technology (Davis, 1989)—and extends its applicability to the specific context of GenAI for OBE-based curriculum development. Lecturers who believe that GenAI helps them complete curriculum development tasks more efficiently, produce higher-quality curriculum documents, achieve better constructive alignment, and make more informed curriculum decisions are significantly more likely to intend to use the technology in their professional practice.

The strength of this relationship is consistent with meta-analytic evidence from broader TAM research, which has consistently identified perceived usefulness as the strongest predictor of behavioral intention across diverse technologies and user populations (Scherer et al., 2019). Within the specific domain of GenAI in education, the finding aligns with Habibi et al. (2023), who reported perceived usefulness as a significant predictor of ChatGPT acceptance among Indonesian university students, and with Shahzad et al. (2025), who confirmed the centrality of usefulness beliefs in GenAI adoption among higher education stakeholders. The present study extends these findings by demonstrating that perceived usefulness remains a robust predictor when operationalized specifically for OBE-related tasks—outcome formulation, alignment mapping, assessment design, and curriculum evaluation—rather than for general academic purposes.

The task-specific nature of the perceived usefulness construct in this study is theoretically significant. The items specifically referenced OBE tasks: efficiency in completing "OBE-based curriculum development tasks," quality improvement in "curriculum documents," alignment of "learning outcomes, learning activities, and assessment methods," and decision-making in "curriculum development." This specificity enhances the construct's predictive validity within the defined context and supports the argument that technology acceptance research benefits from domain-specific operationalization rather than generic usefulness measures that may obscure variation across different professional tasks (Venkatesh & Davis, 2000). The finding that perceived usefulness for OBE tasks specifically drives acceptance suggests that institutional strategies to promote GenAI adoption should demonstrate the technology's utility for concrete curriculum development tasks—showing lecturers precisely how GenAI can assist with outcome formulation, alignment verification, or rubric development—rather than promoting its benefits in abstract terms.

The importance-performance map analysis further illuminated the strategic position of perceived usefulness. With the highest importance score among the direct predictors of acceptance (total effect = 0.353) and a relatively high performance score (69.75), perceived usefulness represents an area where current perceptions are already favorable but where further enhancement would yield meaningful gains in acceptance. This suggests that institutions should both

maintain current usefulness perceptions through continued demonstration of GenAI's practical value and seek to extend these perceptions to curriculum tasks where lecturers may not yet fully appreciate AI's potential contribution.

4. Trust as a Parallel Pathway to Acceptance

The fifth hypothesis, which posited a positive effect of trust on acceptance, was strongly supported ($\beta = 0.328$, $p < 0.001$, $f^2 = 0.156$). Trust in GenAI—belief in its capacity to provide useful initial recommendations, its reliability when outputs are verified, its competence in curriculum-support tasks, and confidence in using it under appropriate guidelines—emerged as a significant predictor of lecturers' behavioral intention, comparable in magnitude to the effect of perceived usefulness. This finding affirms that acceptance of GenAI for curriculum development is not solely a matter of cognitive utility assessment; it also involves a relational dimension of confidence in the technology's reliability and the conditions for its responsible use.

This finding contributes to the emerging consensus that trust is an essential construct in understanding AI acceptance in educational settings. Choi et al. (2023) demonstrated that pedagogical beliefs and perceived trust significantly influenced teachers' acceptance of AI tools, arguing that the relational dimension of human-AI interaction has been undertheorized in traditional TAM research. Balaskas et al. (2025) found that trust mediated the effects of perceived usefulness and perceived ease of use on ChatGPT adoption intention, while perceived risk exerted a parallel negative mediating effect. Al-Abdullatif, (2024) reported that perceived trust mediated the relationship between Intelligent TPACK and GenAI acceptance. The present study extends this line of inquiry by positioning trust as a parallel mediator alongside perceived usefulness, testing whether these two mechanisms operate simultaneously and independently in transmitting the influence of AI literacy to acceptance.

The comparable magnitude of the trust effect ($\beta = 0.328$) and the perceived usefulness effect ($\beta = 0.353$) is noteworthy. While the slightly larger coefficient for perceived usefulness is consistent with TAM's theoretical emphasis on instrumental utility as the primary driver of technology acceptance, the non-trivial contribution of trust indicates that models excluding trust may substantially underestimate the determinants of GenAI acceptance in professional contexts where the stakes of technology use are high. In curriculum development, where AI errors could compromise program coherence or assessment validity, trust is not merely a peripheral consideration but a central psychological prerequisite for adoption. Lecturers must believe not only that GenAI can be useful but also that it can be relied upon within appropriate boundaries.

The importance-performance map analysis revealed that trust recorded a lower performance score (65.25) than perceived usefulness (69.75), indicating that lecturers currently perceive lower levels of trust in GenAI than perceived usefulness. This trust deficit represents a strategic opportunity for institutional intervention. Building trust may require approaches distinct from those that build perceived usefulness. While usefulness perceptions can be enhanced through demonstrations of GenAI's capabilities and workflow integration, trust may require additional conditions: transparent communication about GenAI's limitations and error rates, establishment of clear institutional policies governing AI use in curriculum work, provision of verification protocols and quality assurance checkpoints, and accumulation of positive personal experience through supervised practice. The finding that trust in this study was operationalized as calibrated trust—explicitly incorporating verification

conditions—suggests that institutions should not pursue unconditional trust in AI but should instead foster the conditions for confident use within known boundaries.

5. The Direct and Indirect Pathways from AI Literacy to Acceptance

The third hypothesis, which posited a direct positive effect of AI literacy on acceptance, was supported ($\beta = 0.238$, $p < 0.001$), though the effect size was notably smaller than the mediated pathways. This finding indicates that AI literacy influences lecturers' behavioral intention to use GenAI partially through its own direct channel—independent of perceived usefulness and trust. Lecturers who are more AI-literate may be more willing to adopt GenAI for curriculum development simply because they feel competent to do so, regardless of their current assessment of the technology's usefulness or their degree of trust.

This direct effect is consistent with self-efficacy theory, which suggests that capability beliefs directly influence behavioral intentions in addition to any indirect effects through outcome expectations (Bandura, 1997). In the context of GenAI adoption, AI literacy functions as a form of domain-specific self-efficacy: the belief that one can successfully interact with AI systems to accomplish curriculum development tasks. This self-efficacy belief may motivate adoption through increased comfort with the technology, reduced anxiety about making errors, and greater confidence in one's ability to manage the technology's limitations. The finding aligns with Al-Abdullatif (2024), who reported a significant direct relationship between AI literacy and GenAI acceptance, though the present study's parallel mediation model provides a more nuanced understanding of how this direct effect operates alongside indirect pathways.

However, the substantially larger magnitude of the total indirect effect ($\beta = 0.371$) compared to the direct effect ($\beta = 0.238$) is the more theoretically significant finding. The mediation analysis revealed that 65.5% of AI literacy's total effect on acceptance operates through the combined mediating mechanisms of perceived usefulness and trust. This pattern of complementary partial mediation indicates that while AI literacy does influence acceptance directly, the primary mechanisms through which AI competence translates into adoption intention are cognitive and relational: AI literacy enables lecturers to recognize GenAI's usefulness for their specific professional tasks, and it enables them to develop calibrated trust in the technology's reliability and appropriate use conditions. In other words, AI literacy matters for acceptance largely because it shapes what lecturers believe about GenAI—both in terms of its instrumental value and its trustworthiness—rather than simply because it makes them feel competent to use it.

This finding carries significant theoretical implications for TAM research in AI contexts. It suggests that AI literacy should be conceptualized not merely as an external variable with direct effects on acceptance but as a foundational competence that activates the cognitive and relational mechanisms central to TAM. Theoretically, this positions AI literacy as an enabling condition for the formation of the beliefs—perceived usefulness and trust—that directly drive behavioral intention. Practically, it implies that institutional interventions to promote GenAI adoption should prioritize AI literacy development not as an end in itself but as a means to enable the formation of positive usefulness perceptions and calibrated trust.

6. Parallel Mediation: Comparing the Utility and Trust Pathways

The parallel mediation model tested two distinct mediating mechanisms—a cognitive utility pathway (AI literacy $\rightarrow$ perceived usefulness $\rightarrow$ acceptance) and

a relational confidence pathway (AI literacy → trust → acceptance)—and found both to be significant. The indirect effect through perceived usefulness ($\beta = 0.198$, VAF = 35.0%) was slightly larger than the indirect effect through trust ($\beta = 0.173$, VAF = 30.6%), though the overlapping confidence intervals suggest that this difference is not statistically significant. This pattern indicates that perceived usefulness and trust function as comparably important mediating mechanisms, with neither pathway clearly dominating the other.

The finding that both pathways are significant and roughly comparable in magnitude has several theoretical implications. First, it validates the theoretical distinction between cognitive utility assessment and relational confidence as distinct psychological mechanisms in technology acceptance. While perceived usefulness captures an instrumental evaluation—"will this technology help me perform better?"—trust captures a relational evaluation—"can I rely on this technology within appropriate boundaries?" The significant independent contribution of each pathway, even when the other is controlled, demonstrates that these are not redundant constructs but distinct mechanisms that both transmit the influence of AI literacy to acceptance.

Second, the comparable magnitude of the two pathways suggests that models examining only one mechanism may provide an incomplete picture of how AI literacy translates into acceptance. Al-Abdullatif (2024) examined a serial mediation model where AI literacy influenced acceptance through Intelligent TPACK and perceived trust, with trust positioned as the sole mediating mechanism. The present study's parallel mediation findings suggest that adding perceived usefulness as a parallel mediator captures an additional and equally important pathway. Similarly, studies focusing primarily on perceived usefulness as the mediator of external variable effects (as in traditional TAM research) may underestimate the relational dimension of technology acceptance, particularly for AI systems where trust concerns are salient.

Third, the finding has practical implications for intervention design. The comparable mediation effects suggest that institutional strategies to promote GenAI adoption should address both pathways simultaneously. Efforts to demonstrate GenAI's utility for specific OBE tasks—showing lecturers how AI can accelerate outcome formulation, improve alignment verification, or enhance assessment design—will activate the cognitive utility pathway. Simultaneously, efforts to build calibrated trust—through transparent communication about limitations, provision of verification protocols, establishment of ethical guidelines, and opportunities for supervised practice—will activate the relational confidence pathway. Interventions that address only one pathway may achieve suboptimal results by leaving the other mechanism underdeveloped.

7. Contextual Considerations: Islamic Higher Education

While the study did not hypothesize context-specific effects, the findings merit discussion within the specific institutional context of Indonesian Islamic higher education. The sample comprised exclusively lecturers from State Islamic Universities, State Islamic Institutes, State Islamic Colleges, and private Islamic higher education institutions—institutions that navigate the dual mandate of meeting national OBE standards while maintaining distinctive Islamic scholarly traditions and integrating Islamic values with contemporary knowledge.

The moderate to substantial explanatory power of the model ($R^2 = 0.472$ for acceptance) within this context suggests that the TAM-based framework, extended with AI literacy and trust, generalizes well to Islamic higher education settings. The core psychological mechanisms—utility assessment and calibrated trust—appear to operate similarly across institutional types, disciplinary

backgrounds, and demographic categories within the sample. This finding supports the cross-contextual robustness of the extended TAM framework while also highlighting the importance of contextualizing measurement instruments to the specific professional tasks and institutional settings under study.

Nevertheless, several contextual observations warrant consideration. First, the high percentage of respondents with OBE training experience (78.2%) reflects the substantial investment Indonesian Islamic higher education institutions have made in OBE capacity building, driven by national accreditation requirements. This context of relatively widespread OBE familiarity may have influenced the sample's capacity to meaningfully assess GenAI's usefulness for OBE-specific tasks—a capacity that may be lower in contexts where OBE implementation is less mature. Second, the finding that AI literacy's evaluation dimension recorded the lowest mean ($M = 3.38$) may reflect broader patterns in Indonesian higher education regarding critical digital literacy. While basic operational skills have been developed through general technology use, the critical evaluation of AI-generated academic content may represent a newer competency that has not yet been systematically developed through professional development programs. Third, the ethical dimension of AI literacy ($M = 3.64$) suggests moderate confidence in understanding and applying ethical principles. In Islamic higher education contexts, where knowledge curation carries particular cultural and religious significance, ethical considerations regarding AI's role in curriculum development may extend beyond standard privacy and attribution concerns to encompass questions about the relationship between technological assistance and scholarly authority in Islamic educational traditions.

These contextual considerations suggest directions for future research, including comparative studies across different institutional types within Indonesia, cross-national comparisons with Islamic higher education in other countries, and qualitative investigations into the specific ethical and cultural considerations that shape AI acceptance in Islamic educational settings.

8. Theoretical Implications

This study makes several distinct theoretical contributions to the literature on technology acceptance, AI literacy, and GenAI in higher education.

First, the study extends the Technology Acceptance Model by positioning multidimensional AI literacy as an antecedent that influences acceptance through both direct and indirect pathways. While TAM has been extended with various external variables in prior research, the specific positioning of AI literacy as a higher-order construct comprising awareness, usage, evaluation, and ethics dimensions—and the demonstration that this construct operates through both cognitive and relational mediators—represents a novel theoretical contribution. The finding that 65.5% of AI literacy's total effect on acceptance is mediated through perceived usefulness and trust provides empirical support for the proposition that AI literacy functions primarily as an enabling condition that activates the belief mechanisms central to TAM, rather than as a direct determinant of intention independent of these beliefs.

Second, the study introduces and empirically validates a parallel mediation model that simultaneously tests cognitive utility and relational confidence as distinct mediating mechanisms. This model configuration advances beyond prior research in several ways. Unlike Al-Abdullatif (2024), who examined serial mediation where AI literacy → Intelligent TPACK → perceived trust → acceptance, the present model tests whether perceived usefulness and trust operate as parallel rather than sequential mediators—a theoretically distinct proposition with different practical implications. Unlike Balaskas et al. (2025),

who examined trust and risk as parallel mediators of TAM belief effects, the present model positions AI literacy as the exogenous antecedent, shifting the theoretical focus from general technology beliefs to specific AI competence. The finding that both mediating pathways are significant and comparable in magnitude provides empirical justification for models that incorporate both cognitive and relational mechanisms, rather than privileging one over the other.

Third, the study extends the emerging literature on GenAI acceptance into the underexplored domain of curriculum development. Prior research has predominantly examined GenAI acceptance for general teaching and learning purposes (Habibi et al., 2023; Shahzad et al., 2025), lesson planning (Acquah et al., 2024), or writing support. The present study's focus on OBE-based curriculum development—a complex professional task involving outcome formulation, constructive alignment, assessment design, and continuous improvement—demonstrates that task-specific acceptance models can provide nuanced insights beyond those yielded by general acceptance studies. The finding that perceived usefulness for specific OBE tasks drives acceptance supports the theoretical argument that technology acceptance research benefits from domain-specific operationalization of constructs (Venkatesh & Davis, 2000).

Fourth, the study contributes to the operationalization and measurement of AI literacy in professional educational contexts. By adapting the four-dimensional AI literacy framework (awareness, usage, evaluation, ethics) from Wang et al. (2023) and Ng et al. (2021) to the specific context of curriculum development, and by validating this adaptation through rigorous measurement model assessment, the study provides an instrument that can be adapted for future research on AI literacy in other professional educational domains. The finding that the four dimensions function as reflective indicators of a higher-order AI literacy construct supports the theoretical conceptualization of AI literacy as an integrated professional competency rather than a set of independent skills.

9. Practical Implications

The findings of this study carry actionable implications for Islamic higher education institutions, curriculum development units, and individual lecturers seeking to promote responsible GenAI adoption for OBE-based curriculum development.

First, AI literacy development should be a strategic institutional priority. The finding that AI literacy is the foundational antecedent driving both perceived usefulness and trust—and through them, acceptance—indicates that investments in AI literacy training will yield multiplicative benefits. Training programs should be comprehensive, addressing all four dimensions of AI literacy: awareness of capabilities and limitations, practical usage skills including prompt formulation and output refinement, critical evaluation of AI-generated content for accuracy and alignment, and ethical understanding encompassing data privacy, disclosure, and professional boundaries. The finding that the evaluation dimension recorded the lowest mean score suggests that particular emphasis should be placed on developing lecturers' capacity to critically assess AI-generated learning outcomes for measurability, to verify the accuracy of AI-generated curriculum recommendations, and to evaluate the alignment between AI-generated outcomes, activities, and assessments. Training that neglects evaluation in favor of basic operational skills will produce lecturers who can use GenAI but cannot use it well or responsibly.

Second, institutions should address both the utility and trust pathways to maximize adoption. The comparable mediation effects of perceived usefulness

and trust suggest that interventions focusing exclusively on demonstrating AI's usefulness—without addressing trust concerns—will achieve suboptimal results, and vice versa. Utility-focused strategies might include developing libraries of model prompts for common OBE tasks (outcome formulation, mapping analysis, rubric development), showcasing concrete examples of AI-assisted curriculum development with demonstrable quality improvements, and integrating GenAI tools into existing curriculum development workflows and templates. Trust-focused strategies might include establishing clear institutional policies governing AI use in curriculum work, developing verification protocols and quality assurance checkpoints for AI-generated content, transparently communicating both the capabilities and limitations of GenAI tools, providing supervised practice opportunities where lecturers can build positive experience with AI in low-stakes contexts, and ensuring that institutional data protection policies address the specific risks of using public AI platforms for curriculum materials.

Third, the human-in-the-loop principle should be institutionalized. The conceptualization of trust in this study as calibrated trust—contingent on verification—and the operationalization of GenAI as a decision-support tool rather than an autonomous designer both point toward the necessity of formalizing human oversight in AI-assisted curriculum development. Institutions should establish protocols requiring human verification of all AI-generated curriculum content before adoption, designate responsibility for verification at appropriate levels (individual lecturers for course-level materials, curriculum committees for program-level documents), and develop criteria for assessing whether AI-generated outputs meet institutional standards for quality, alignment, and appropriateness. These protocols should be designed to enable confident use of GenAI within known boundaries rather than to discourage use entirely.

Fourth, ethical frameworks for AI use in curriculum development should be developed and communicated. The moderate mean score for the ethics dimension of AI literacy ($M = 3.64$) indicates that while lecturers have some ethical awareness, there is room for more systematic institutional guidance. Ethical frameworks should address: data privacy and the prohibition of entering confidential student or institutional data into public AI platforms; disclosure and attribution practices for AI assistance in curriculum documents; boundaries of AI's appropriate role in academic decision-making, particularly for decisions requiring professional pedagogical judgment; and considerations specific to Islamic higher education contexts, such as the relationship between AI-generated content and scholarly authority in Islamic knowledge traditions. UNESCO's guidance on GenAI in education (Miao & Holmes, 2023) and the AI Competency Framework for Teachers (Miao & Cukurova, 2024) provide internationally recognized starting points that can be adapted to local institutional contexts.

Fifth, professional development should be differentiated based on current AI literacy levels. The descriptive statistics revealed variation across AI literacy dimensions and across respondents, suggesting that a one-size-fits-all training approach may be suboptimal. Institutions might consider diagnostic assessments of lecturers' AI literacy profiles to identify specific development needs: lecturers strong in awareness but weak in evaluation may need focused training on verification strategies; lecturers strong in usage but weak in ethics may need guidance on data privacy and disclosure practices; lecturers with uniformly low AI literacy may need comprehensive foundational programs before task-specific application training. The importance-performance map analysis identified AI literacy as the construct with the highest total importance for acceptance but the

lowest performance score, reinforcing the prioritization of AI literacy development as the intervention with the greatest potential return.

Sixth, curriculum development units and quality assurance offices should develop AI-integrated workflows. Beyond individual lecturer adoption, institutions can facilitate GenAI integration by embedding AI tools and protocols into formal curriculum development and review processes. This might include: developing institutional prompt libraries for common curriculum tasks aligned with national higher education standards and accreditation criteria; creating templates that combine AI-generated drafts with required human verification checkpoints; establishing peer review processes for AI-assisted curriculum materials; and documenting cases of effective AI-assisted curriculum development to serve as models and build institutional knowledge.

Finally, institutions should monitor and evaluate the impact of GenAI integration on curriculum quality. While this study focused on acceptance as behavioral intention, the ultimate goal of GenAI adoption is improved curriculum quality and enhanced student learning. Institutions should develop mechanisms to assess whether AI-assisted curriculum development actually yields improvements in outcome clarity, constructive alignment, assessment validity, and ultimately, student achievement of learning outcomes. Such evidence would not only justify continued investment in AI integration but would also inform the refinement of AI use protocols and training programs.

CONCLUSION

This study examined how AI literacy, perceived usefulness, and trust shape lecturers' acceptance of generative AI for OBE-based curriculum development in Indonesian Islamic higher education. Using an extended Technology Acceptance Model with a parallel mediation framework, the findings show that AI literacy significantly enhances perceived usefulness ($\beta = 0.562$) and trust ($\beta = 0.527$), which subsequently mediate lecturers' behavioral intention to use generative AI. Perceived usefulness and trust account for 35.0% and 30.6% of the total effect of AI literacy on acceptance, respectively, while AI literacy also exerts a smaller but significant direct effect ($\beta = 0.238$). Overall, the model explains 47.2% of the variance in lecturers' behavioral intention. These results indicate that AI literacy functions as a foundational professional competence whose influence operates through two distinct but complementary pathways: an instrumental assessment of whether generative AI improves the efficiency and quality of OBE-related tasks, and a trust-based assessment of whether the technology can be used reliably within appropriate verification boundaries.

The findings extend the Technology Acceptance Model by demonstrating that perceived usefulness and trust are not redundant mechanisms but equally important determinants of generative AI acceptance in a task-specific professional context. Perceived usefulness reflects lecturers' evaluation of the contribution of generative AI to activities such as formulating learning outcomes, mapping curriculum components, establishing constructive alignment, and designing assessments. Trust reflects confidence that generative AI can provide relevant and sufficiently reliable assistance when its outputs are critically evaluated by human experts. For Islamic higher education institutions, these findings suggest that responsible adoption should not be limited to providing technological access or general training. Institutions should strengthen lecturers' AI literacy across the dimensions of awareness, usage, evaluation, and ethics, demonstrate the practical value of generative AI through authentic OBE-related activities, and establish clear guidelines, verification procedures, data-protection standards, and human-in-the-loop mechanisms. AI-generated content should

therefore be treated as preliminary academic support, while responsibility for curriculum validity, institutional alignment, academic integrity, and ethical acceptability remains with lecturers, curriculum teams, study programs, and quality assurance units.

The findings should nevertheless be interpreted in light of several limitations. The cross-sectional design limits definitive causal inference, while the use of self-reported perceptions and behavioral intention rather than actual use may introduce common method variance and an intention–behavior gap. Purposive sampling also restricts statistical generalizability, and the model did not examine potential moderators such as institutional support, disciplinary differences, teaching experience, prior OBE training, or specific generative AI platforms. In addition, the study did not assess the actual quality of AI-assisted curriculum documents. Future research should employ longitudinal, experimental, and qualitative designs; incorporate objective measures of AI literacy, actual use, and curriculum quality; and conduct comparisons across disciplines, institutions, and national contexts. In conclusion, AI literacy should be understood not merely as a technical ability but as a foundational professional competence that shapes how lecturers recognize the usefulness of generative AI, develop calibrated trust in its capabilities, and responsibly integrate it into the complex process of OBE-based curriculum development.

REFERENCES

- Acquah, B. Y. S., Arthur, F., Salifu, I., Quayson, E., & Nortey, S. A. (2024). Preservice teachers' behavioural intention to use artificial intelligence in lesson planning: A dual-staged PLS-SEM-ANN approach. *Computers and Education: Artificial Intelligence*, 7, 100307. <https://doi.org/10.1016/j.caeai.2024.100307>
- Al-Abdullatif, A. M. (2024). Modeling teachers' acceptance of generative artificial intelligence use in higher education: The role of AI literacy, Intelligent TPACK, and perceived trust. *Education Sciences*, 14(11), 1209. <https://doi.org/10.3390/educsci14111209>
- Balaskas, S., Tsiantos, V., Chatzifotiou, S., & Rigou, M. (2025). Determinants of ChatGPT adoption intention in higher education: Expanding on TAM with the mediating roles of trust and risk. *Information*, 16(2), 82. <https://doi.org/10.3390/info16020082>
- Bandura, A. (1997). *Self-efficacy: The exercise of control*. W.H. Freeman.
- Becker, J.-M., Klein, K., & Wetzels, M. (2012). Hierarchical latent variable models in PLS-SEM: Guidelines for using reflective-formative type models. *Long Range Planning*, 45(5–6), 359–394. <https://doi.org/10.1016/j.lrp.2012.10.001>
- Biggs, J. (1996). Enhancing teaching through constructive alignment. *Higher Education*, 32(3), 347–364. <https://doi.org/10.1007/BF00138871>
- Chan, C. K. Y. (2023). A comprehensive AI policy education framework for university teaching and learning. *International Journal of Educational Technology in Higher Education*, 20, 38. <https://doi.org/10.1186/s41239-023-00408-3>
- Choi, S., Jang, Y., & Kim, H. (2023). Influence of pedagogical beliefs and perceived trust on teachers' acceptance of educational artificial intelligence tools. *International Journal of Human–Computer Interaction*, 39(4), 910–922. <https://doi.org/10.1080/10447318.2022.2049145>
- Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly: Management Information Systems*, 13(3), 319–339. <https://doi.org/10.2307/249008>
- Farrokhnia, M., Banihashem, S. K., Noroozi, O., & Wals, A. (2024). A SWOT analysis

- of ChatGPT: Implications for educational practice and research. *Innovations in Education and Teaching International*, 61(3), 460–474. <https://doi.org/10.1080/14703297.2023.2195846>
- Habibi, A., Muhaimin, M., Danibao, B. K., Wibowo, Y. G., Wahyuni, S., & Octavia, A. (2023). ChatGPT in higher education learning: Acceptance and use. *Computers and Education: Artificial Intelligence*, 5, 100190. <https://doi.org/10.1016/j.caeai.2023.100190>
- Hair, J. F., Hult, G. T. M., Ringle, C. M., & Sarstedt, M. (2022). *A primer on partial least squares structural equation modeling (PLS-SEM)* (3rd ed.). Sage. <https://doi.org/10.1007/978-3-030-80519-7>
- Harden, R. M. (2007). Outcome-based education: The future is today. *Medical Teacher*, 29(7), 625–629. <https://doi.org/10.1080/01421590701729930>
- Kock, N. (2015). Common method bias in PLS-SEM: A full collinearity assessment approach. *International Journal of E-Collaboration*, 11(4), 1–10. <https://doi.org/10.4018/ijec.2015100101>
- Lo, C. K. (2023). What is the impact of ChatGPT on education? A rapid review of the literature. *Education Sciences*, 13(4), 410. <https://doi.org/10.3390/educsci13040410>
- Mayer, R. C., Davis, J. H., & Schoorman, F. D. (1995). An integrative model of organizational trust. *Academy of Management Review*, 20(3), 709–734. <https://doi.org/10.2307/258792>
- Miao, F., & Cukurova, M. (2024). *AI competency framework for teachers*. UNESCO. <https://doi.org/10.54675/ZJTE2084>
- Miao, F., & Holmes, W. (2023). *Guidance for generative AI in education and research*. UNESCO. <https://doi.org/10.54675/EWZM9535>
- Ng, D. T. K., Leung, J. K. L., Chu, S. K. W., & Qiao, M. S. (2021). Conceptualizing AI literacy: An exploratory review. *Computers and Education: Artificial Intelligence*, 2, 100041. <https://doi.org/10.1016/j.caeai.2021.100041>
- Preacher, K. J., & Hayes, A. F. (2008). Asymptotic and resampling strategies for assessing and comparing indirect effects in multiple mediator models. *Behavior Research Methods*, 40(3), 879–891. <https://doi.org/10.3758/BRM.40.3.879>
- Scherer, R., Siddiq, F., & Tondeur, J. (2019). The technology acceptance model (TAM): A meta-analytic structural equation modeling approach to explaining teachers' adoption of digital technology in education. *Computers & Education*, 128, 13–35. <https://doi.org/10.1016/j.compedu.2018.09.009>
- Shahzad, M. F., Xu, S., & Asif, M. (2025). Factors affecting generative artificial intelligence, such as ChatGPT, use in higher education: An application of technology acceptance model. *British Educational Research Journal*, 51(2), 489–513. <https://doi.org/10.1002/berj.4084>
- Spady, W. G. (1994). *Outcome-based education: Critical issues and answers*. American Association of School Administrators.
- UNESCO. (2023). *Guidance for generative AI in education and research*. UNESCO Publishing.
- Venkatesh, V., & Davis, F. D. (2000). A theoretical extension of the Technology Acceptance Model: Four longitudinal field studies. *Management Science*, 46(2), 186–204. <https://doi.org/10.1287/mnsc.46.2.186.11926>
- Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User acceptance of information technology: Toward a unified view. *MIS Quarterly*, 27(3), 425–478. <https://doi.org/10.2307/30036540>
- Wang, B., Rau, P.-L. P., & Yuan, T. (2023). Measuring user competence in using artificial intelligence: Validity and reliability of artificial intelligence literacy scale. *Behaviour & Information Technology*, 42(9), 1324–1337. <https://doi.org/10.1080/0144929X.2022.2072768>

- Weick, K. E. (1995). *Sensemaking in organizations*. Sage.
- Zhao, X., Lynch, J. G., & Chen, Q. (2010). Reconsidering Baron and Kenny: Myths and truths about mediation analysis. *Journal of Consumer Research*, 37(2), 197–206. <https://doi.org/10.1086/651257>