

**INTERNATIONAL JOURNAL OF
MODERN EDUCATION
(IJMOE)**www.gaexcellence.com/ijmoe**VALIDATING THE ADEM SCALE: A SECOND-ORDER
MEASURE OF STUDENTS' AI DISPOSITION AND
ADAPTATION IN RELATION TO AFFECTIVE AND
IDENTITY OUTCOMES**

Ateerah Abdul Razak^{1,2*}, Nur Hafifah Jamalludin³, Amanina Abdul Razak⁴, Azahah Abu Hassan Shaari⁵

¹Institut Penyelidikan dan Pengurusan Kemiskinan (InsPeK), Universiti Malaysia Kelantan, Malaysia

²Faculty of Language Studies and Human Development, Universiti Malaysia Kelantan, Malaysia

 ateerah@umk.edu.my

 <https://orcid.org/0000-0003-2987-2999>

³Faculty of Language Studies and Human Development, Universiti Malaysia Kelantan, Malaysia

 hafifah.j@umk.edu.my

 <https://orcid.org/0009-0007-1991-571X>

⁴Pusat Asasi, UiTM Cawangan Selangor Kampus Dengkil

 amanina81@uitm.edu.my

 <https://orcid.org/0000-0002-5226-7968>

⁵Faculty of Social Sciences and Humanities, The National University of Malaysia, Bangi, Malaysia

 azahah@ukm.edu.my

 <https://orcid.org/0000-0003-0726-0043>

*Corresponding Author

Article Info:**Article history:**

Received date: 05.07.2026

Revised date: 06.08.2026

Accepted date: 08.09.2026

Published date: 15.09.2026

To cite this document:

Razak, A. A., Jamalludin, N. H., Razak, A. A., Shaari, A. A. H. (2026). Validating The ADEM Scale: A Second-Order Measure of Students' AI Disposition and Adaptation in Relation to Affective and Identity Outcomes. *International Journal of Modern Education*, 8(31), 297-314.

Abstract:

The rapid integration of artificial intelligence (AI) into university learning has created a need for a comprehensive measure that considers not only students' ability to use AI but also their intellectual readiness, emotional regulation, ethical values, and identity as learners. This cross-sectional study assessed the psychometric quality of the 70-item AI Disposition and Adaptation Model (ADEM) and examined how AI Adaptability, Epistemological Awareness, Resource Integration, and Intellectual Readiness relate to students' overall adaptation to AI-supported learning. ADEM was represented by three interconnected dimensions: Emotional Balance, Axiological Alignment, and Ontological Perspective. Survey data were collected from 280 higher-education students. The findings demonstrated strong reliability and convergent validity across the measured dimensions, while also identifying some overlap between closely related constructs that should be considered when interpreting the instrument. AI Adaptability, Epistemological Awareness, Resource Integration, and Intellectual Readiness were significantly associated with ADEM. Intellectual Readiness also partially explained how students' adaptability, critical evaluation of AI-generated knowledge,

and integration of AI with authoritative academic resources contributed to their broader adaptation to AI-supported learning. This study introduces a multidimensional framework that extends beyond technical competence by incorporating emotional well-being, responsible knowledge evaluation, ethical conduct, and learner agency. The ADEM instrument offers a practical basis for assessing students' strengths and development needs, informing AI-literacy curricula, designing targeted educational interventions, and evaluating whether AI-supported learning initiatives promote capable, critical, balanced, and ethically responsible engagement with AI.

Keywords:

AI Adaptability; AI-Supported Learning; Ethical AI Engagement; Higher Education; Intellectual Readiness.

DOI: 10.35631/IJMOE.831018



© The authors (2026). This is an Open Access article distributed under the terms of the Creative Commons Attribution (CC BY NC) (<http://creativecommons.org/licenses/by-nc/4.0/>), which permits non-commercial re-use, distribution, and reproduction in any medium, provided the original work is properly cited. For commercial re-use, please contact ijmoe@gaexcellence.com.

Introduction

Generative-AI tools now mediate routine study tasks such as drafting, debugging, and information searching, reshaping how students plan, evaluate, and integrate sources. While such tools can lower cognitive effort and expand access to feedback, they also raise risks of over-reliance, uncritical acceptance of outputs, and uncertainty about authorship and academic identity. In this study, ethical-identity outcomes refer specifically to values-consistent AI use, represented by Axiological Alignment (AA), and the preservation of learner identity, authenticity, and agency, represented by Ontological Perspective (OP). Institutions therefore require instruments that go beyond usage counts to capture students' disposition and adaptation to AI together with these affective and ethical-identity outcomes (Giannakos et al., 2024).

The ADEM framework comprises seven first-order reflective measures: AI Adaptability (AD), Epistemological Awareness (EA), Resource Integration (RI), Intellectual Readiness (IR), Emotional Balance (EB), Axiological Alignment (AA), and Ontological Perspective (OP). EB, AA, and OP represent the first-order components of ADEM, which is modeled as a second-order reflective construct. In the structural model, AD, EA, and RI are exogenous predictors, IR is a proximal predictor and mediator, and ADEM is the endogenous outcome.

Building on prior ADEM work, this study has two objectives. First, it evaluates the reliability, convergent validity, and discriminant validity of the 70-item instrument using outer loadings, composite reliability (CR), average variance extracted (AVE), and numerical heterotrait–monotrait ratios (HTMT). Second, it tests whether AD, EA, RI, and IR contribute uniquely to the second-order ADEM outcome and whether IR mediates the effects of AD, EA, and RI.

Theoretical Framework and Hypotheses

Grounded in self-efficacy and cognitive load perspectives, students who possess stronger Intellectual Readiness (IR), reflected in clear mental models of what AI can and cannot do, and greater AI Adaptability (AD), expressed in fluent and flexible tool use in authentic tasks, should experience less uncertainty and lower processing strain during AI-supported study. Perceived capability supports effort regulation and affective control (Gkintoni et al., 2025), while fluent skills and well-structured practice reduce working-memory demands (Premeti et al., 2024), together stabilizing the affective climate of learning. Accordingly, IR and AD are expected to relate positively to Emotional Balance (EB) and, more broadly, to the higher-order ADEM outcome that aggregates EB, Axiological Alignment (AA), and Ontological Perspective (OP) through a reflective–reflective specification.

Complementing skills and comfort, Epistemological Awareness (EA) involves evaluating claims, calibrating trust, and justifying beliefs about AI outputs (Sandoval et al., 2016), whereas Resource Integration (RI) concerns triangulating those outputs with textbooks, peer-reviewed literature, and course materials (Akbarighatar, 2024). The constructs are related but distinct: EA concerns evaluation of knowledge claims, while RI concerns integration and verification across sources. Both are tested as predictors of ADEM and as antecedents of IR. Values-consistent conduct is represented by Axiological Alignment (AA), and coherent learner identity and agency are represented by Ontological Perspective (OP).

Integrating these strands, the measurement framework comprises seven first-order reflective constructs: AD, EA, RI, IR, EB, AA, and OP. EB, AA, and OP are the affective and ethical-identity components of the second-order reflective ADEM outcome. In the structural model, AD, EA, and RI are exogenous predictors, IR is both a proximal predictor and mediator, and ADEM is endogenous. Six hypotheses therefore test three direct effects and three indirect effects through IR.

Directional Hypotheses Aligned With The Reported Structural Model

- H1: AI Adaptability (AD) → ADEM (+)
- H2: Epistemological Awareness (EA) → ADEM (+)
- H3: Resource Integration (RI) → ADEM (+)
- H4: AD → Intellectual Readiness (IR) → ADEM (partial mediation)
- H5: EA → Intellectual Readiness (IR) → ADEM (partial mediation)
- H6: RI → Intellectual Readiness (IR) → ADEM (partial mediation)

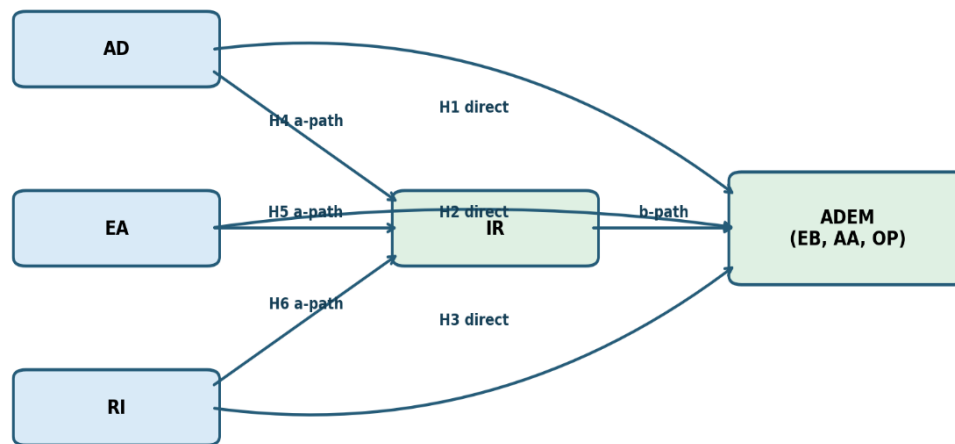


Figure 1: Structural Model and Hypotheses Tested in the Present Study (H1-H6)

Source: Authors' conceptualisation.

ADEM is modeled as a second-order reflective construct represented by Emotional Balance (EB), Axiological Alignment (AA), and Ontological Perspective (OP). H1–H3 specify the direct effects of AD, EA, and RI on ADEM, whereas H4–H6 specify their indirect effects through IR. IR also serves as a proximal predictor of ADEM in each mediation model (Chin et al., 2003).

The second-order specification clarifies how EB, AA, and OP represent a common ADEM outcome while remaining conceptually distinct. EB reflects affective regulation during AI use; AA reflects values-consistent and ethically aligned AI use; and OP reflects learner identity, authenticity, and agency. Modeling their shared variance reduces the complexity of testing three closely related outcomes separately (Diamantopoulos, 2011). H1–H3 test the direct associations of AD, EA, and RI with ADEM, while IR is included as a proximal predictor and mediator in H4–H6.

The mediation pathways provide the mechanism account. H4–H6 propose that IR functions as a proximal cognitive resource through which AD, EA, and RI exert additional influence on ADEM. Mediation is evaluated using each indirect effect ($X \rightarrow IR \rightarrow ADEM$) together with its residual direct effect ($X \rightarrow ADEM$). A significant indirect effect accompanied by a significant direct effect supports partial mediation (Hayes, 2015).

Research Methodology

Design & Participants

A cross-sectional survey was administered to higher-education students ($N = 280$) using a stratified random sampling approach to ensure representation across major strata of the university, including faculties, discipline clusters, and programme levels (Kurniawan et al., 2024). Within each stratum, students were randomly invited to participate if they were currently

enrolled and using AI in their learning processes such as AI search, coding assistants, and retrieval-augmented writing. The instrument was delivered online; participation was voluntary with informed consent, and responses were anonymous. This design yields a cohort reflecting the diversity of AI users in the university context while minimizing selection bias and supporting generalizability of the ADEM measurement and structural findings to the broader higher-education population (Xiao et al., 2023).

Ethics

The protocol received approval from University of Malaysia Kelantan (Approval ID [UMK.A16.800-6/1/1 (12)], see Appendix). Participation was voluntary with informed consent, and responses were anonymous (Hoft, 2021). Inclusion criteria required current enrolment and active AI use in learning, such as AI search, coding assistants, and retrieval-augmented writing; no personally identifying information was collected beyond basic demographics. Data were stored securely in accordance with institutional policy.

Instrument & Scoring

The survey comprised seven reflective constructs, each measured with 10 Likert-type items on a 1–6 scale, with higher scores indicating greater levels of the attribute. Items were grouped by construct; Intellectual Readiness (IR), AI Adaptability (AD), Epistemological Awareness (EA), Resource Integration (RI), Emotional Balance (EB), Axiological Alignment (AA), and Ontological Perspective (OP), and treated as reflective indicators of their respective latent variables. Prior to aggregation, responses were screened for valid ranges and reverse-keyed items were reoriented so that larger values consistently indicated stronger endorsement (Ilhan et al., 2024). Construct scores were computed as the arithmetic mean of available items, with a minimum completion rule of $\geq 50\%$ per construct, at least 5 of 10 items. Records failing this threshold for a given construct were set to missing for that scale. This scoring approach yields interpretable 1–6 composites while preserving sample size by tolerating limited item-level missingness completed (Martinez & Bartholomew, 2017).

Measurement Modelling

The analysis followed a two-stage hierarchical measurement framework. First, all seven first-order constructs (IR, AD, EA, RI, EB, AA, and OP) were assessed using outer loadings, CR, AVE, and HTMT. Second, ADEM was specified as a reflective higher-order outcome represented by EB, AA, and OP. For reproducibility from the supplied dataset, the reported structural and mediation estimates use unit-weighted construct scores, HC3 robust standard errors, R^2 , 5,000-resample bootstrap confidence intervals, multicollinearity checks, and influence diagnostics (Akter et al., 2017; Hair & Alamer, 2022).

Stage 1 involved estimating reflective measurement models for EB, AA, OP, IR, AD, EA, and RI using Mode-A (correlation weighting) to obtain latent variable scores and assess indicator quality (Hair et al., 2021). Diagnostics included outer loadings with a target of ≥ 0.70 , while indicators in the 0.40–0.70 band were retained only when theoretically essential and when $CR \geq 0.70$ and $AVE \geq 0.50$ remained adequate. Composite reliability ($CR \geq 0.70$) and average variance extracted ($AVE \geq 0.50$) were used to establish convergent validity, and HTMT was applied to assess discriminant validity, with values preferably below 0.85 and acceptable up to

0.90, and with bootstrapped HTMT confidence intervals inspected where feasible (Méndez-Suárez, 2021).

Stage 2 specified ADEM as a second-order reflective construct represented by the Stage-1 scores for EB, AA, and OP. The estimated second-order loadings were .898, .920, and .936, respectively; second-order CR was .941 and AVE was .843. This specification captures shared variance in affective balance, values-consistent conduct, and learner identity while preserving the conceptual distinction among these components (Becker et al., 2012; Duarte & Amaro, 2018).

Structural Tests

The structural model specified ADEM as the endogenous second-order construct, with AD, EA, RI, and IR as simultaneous predictors ($ADEM \sim AD + EA + RI + IR$). Construct scores were calculated as arithmetic means of their ten indicators, and ADEM was calculated as the mean of EB, AA, and OP scores. Paths were estimated using ordinary least squares with HC3 heteroskedasticity-robust standard errors; 5,000 case-resampling bootstraps produced percentile 95% confidence intervals. This score-based implementation is reported transparently so that the estimates are reproducible from the supplied data.

Diagnostics & Robustness

Model assumptions and robustness were evaluated as follows:

Table 1: Model Assumptions and Robustness Checks

Diagnostic Focus	Evaluation Method(s)	Criteria / Outcome
Residual behavior	Fitted-vs-residual plots; Q–Q plots	Pattern-free and approximately normal
Heteroskedasticity	Breusch–Pagan test; White test	Addressed in inference by HC3 robust standard errors
Multicollinearity	Inner-model VIF	VIF = 3.50–5.24; IR = 5.24 flagged as a caution
Influence	Cook’s distance (4/N heuristic)	No cases exceeded 4/N; re-estimation excluding top-influence cases showed no change in signs, magnitudes, or significance

Source: Authors' analysis.

The diagnostics summarized in Table 1 were used to assess the structural estimates. Residual and influence checks did not indicate a change in the substantive conclusions. Predictor VIFs ranged from 3.50 to 5.24; the IR value (5.24) was slightly above the common threshold of 5 and is therefore treated as a caution rather than as evidence of an unqualified model fit (Hickey et al., 2018; Thompson et al., 2017). HC3 standard errors were retained for inference.

Mediation

Indirect effects were tested for $AD \rightarrow IR \rightarrow ADEM$ (H4), $EA \rightarrow IR \rightarrow ADEM$ (H5), and $RI \rightarrow IR \rightarrow ADEM$ (H6) using the product-of-coefficients approach with 5,000 case-resampling

bootstraps and percentile 95% confidence intervals. Each model estimated the a path, b path, indirect effect $a \times b$, and residual direct effect c' . Partial mediation was concluded only when both the indirect and corresponding direct effects were significant (Tofighi, 2020).

Results

Measurement Model (Reflective)

Table 2: Measurement Model: Outer Loadings, CR, and AVE (N = 280)

Construct	Outer loadings (Items 1–10)	CR	AVE
EB	.762, .654, .738, .787, .828, .686, .770, .728, .822, .692	.927	.561
IR	.819, .843, .874, .711, .821, .854, .846, .816, .834, .813	.955	.679
AD	.807, .869, .872, .856, .841, .849, .833, .873, .821, .840	.962	.716
EA	.877, .870, .859, .882, .840, .870, .857, .829, .889, .840	.966	.742
OP	.823, .867, .817, .793, .851, .813, .827, .846, .844, .853	.958	.695
RI	.873, .870, .894, .901, .880, .893, .869, .908, .884, .889	.973	.785
AA	.876, .916, .868, .842, .926, .881, .908, .911, .901, .915	.976	.801

Source: Authors' analysis.

Table 4.1 reports every outer loading in appendix item order. Loadings ranged from .654 to .926; EB2 (.654), EB7 (.686), and EB10 (.692) were below .70 but were retained because the EB scale remained reliable (CR = .927) and convergent validity remained acceptable (AVE = .561). All other constructs showed $CR \geq .955$ and $AVE \geq .679$ (Peterson et al., 2020).

Table 3: Numerical HTMT Matrix for First-Order Constructs

Construct	EB	IR	AD	EA	OP	RI	AA
EB	—						
IR	.923	—					
AD	.885	.919	—				
EA	.746	.837	.790	—			
OP	.814	.856	.881	.848	—		
RI	.695	.788	.758	.858	.844	—	
AA	.759	.791	.803	.813	.853	.814	—

Source: Authors' analysis.

The numerical HTMT matrix shows that 19 of 21 construct pairs were below .90. EB–IR (.923) and IR–AD (.919) exceeded the .90 guideline; AD–EB (.885) also exceeded the stricter .85 guideline. Accordingly, the results support reliability and convergent validity but provide only qualified evidence of discriminant validity. These elevated ratios are reported as a measurement limitation requiring item review and independent replication, not concealed through a blanket validity claim (Rönkkö & Cho, 2020).

The overlap is theoretically plausible because emotional confidence, intellectual readiness, and adaptability may co-develop during AI-supported learning. Nevertheless, conceptual plausibility does not remove the empirical concern. Future validation should examine cross-loadings, revise overlapping items, and test whether the distinctions are stable across institutions and disciplines.

The seven constructs are therefore retained provisionally for structural analysis, with the elevated HTMT ratios carried forward as a limitation. ADEM remains the second-order outcome represented by EB, AA, and OP.

Second-Order Measurement And Structural Model

Table 4: Second-Order ADEM Measurement Results

Component	Loading	95% bootstrap CI
EB	.898	[.863, .924]
AA	.920	[.897, .940]
OP	.936	[.917, .951]
ADEM reliability	CR = .941	AVE = .843

Source: Authors' analysis.

Table 5: Structural Paths to ADEM (HC3 Standard Errors; 5,000 Bootstraps)

Hypothesis / predictor	Path	b	SE (HC3)	p	95% bootstrap CI	Decision
IR (covariate)	IR → ADEM	.233	.047	< .001	[.136, .326]	Significant
H1	AD → ADEM	.340	.063	< .001	[.220, .445]	Supported
H2	EA → ADEM	.154	.072	.033	[.028, .282]	Supported
H3	RI → ADEM	.187	.063	.003	[.076, .309]	Supported

Source: Authors' analysis.

Table 4.4 shows that all four simultaneous predictors were positively associated with ADEM. AD had the largest coefficient ($b = .340$), followed by IR ($b = .233$), RI ($b = .187$), and EA ($b = .154$). H1–H3 were supported. Because the evidence is cross-sectional, these estimates are interpreted as conditional associations rather than causal effects.

Table 6: Model Summary (ADEM as Endogenous Construct)

Endogenous construct	R ²	Adj. R ²	N	Bootstrap resamples
ADEM (2nd order)	.872	.871	280	5,000

Source: Authors' analysis.

The model explained a substantial proportion of variance in ADEM ($R^2 = .872$; adjusted $R^2 = .871$). This variance concerns the shared EB, AA, and OP outcome space and is specific to the present sample, score construction, and structural specification.

Mediation (5,000 resamples; IR as mediator)

Table 7: Mediation Results (IR as Mediator)

Hypothesis	a	b	Indirect	95% CI	Direct c'	Direct 95% CI	N
H4: AD → IR → ADEM	.804	.429	.345	[.243, .468]	.431	[.294, .541]	280

Hypothesis	a	b	Indirect	95% CI	Direct c'	Direct 95% CI	N
H5: EA → IR → ADEM	.764	.560	.428	[.324, .547]	.333	[.194, .448]	280
H6: RI → IR → ADEM	.688	.574	.395	[.306, .491]	.317	[.203, .424]	280

Source: Authors' analysis.

Bootstrapped analyses indicated significant indirect effects through IR for AD (.345, 95% CI [.243, .468]), EA (.428, 95% CI [.324, .547]), and RI (.395, 95% CI [.306, .491]). Their residual direct effects also remained significant; therefore, H4–H6 were supported as partial mediation rather than full mediation.

The results support a dual-pathway interpretation in which adaptability, epistemological awareness, and resource integration are associated with ADEM both directly and indirectly through intellectual readiness. These pathways are statistical mechanisms within cross-sectional data and do not establish temporal causation.

Discussion

The results support complementary capability, epistemic, and resource-practice mechanisms. AD and IR may reduce uncertainty and processing strain; EA supports careful evaluation of AI-generated knowledge; and RI represents the triangulation of AI outputs with authoritative academic sources. Each construct showed a unique positive association with the second-order ADEM outcome after adjustment for the other predictors.

The partial mediation findings indicate that IR functions as a proximal cognitive resource through which AD, EA, and RI are additionally associated with ADEM. Because the residual direct effects remained significant, none of these relationships operated solely through IR. Ethical-identity outcomes refer specifically to AA and OP within the second-order ADEM construct and not to a separate measured variable.

For teaching and curriculum, the findings support a concise paired strategy: build AI capability and intellectual readiness while requiring evidence evaluation and triangulation with textbooks, peer-reviewed literature, and course materials. Claim–evidence prompts, sourcing checklists, and comparison tasks can develop EA and RI without treating confident AI use as sufficient evidence of responsible use (ACRL, 2016; Ecker et al., 2022).

For programme evaluation and policy, the ADEM subscales (EB, AA, OP) can function as pre–post indicators for AI-literacy initiatives at course, programme, and faculty levels. Institutions can integrate these metrics into quality-assurance dashboards to monitor cohort trends in affective well-being and integrity-aligned use, set progress benchmarks, and identify programmes needing targeted support (Susnjak et al., 2022). Such monitoring aligns teaching practice with institutional commitments to human-centred, ethical AI in education and provides actionable feedback loops for continuous improvement.

Limitations and future work. The cross-sectional, self-report, single-context design limits causal inference and generalizability. The evidence is sample-specific, and the high explained variance may partly reflect related self-report measures. Two HTMT ratios exceeded .90 (EB–

IR and IR–AD), AD–EB exceeded .85, and the IR VIF was slightly above 5; these findings indicate construct overlap that should be examined through item refinement, cross-loadings, confirmatory factor analysis, and independent replication. Multi-site and longitudinal studies should test measurement invariance and temporal ordering before stronger validation or causal claims are made.

Conclusion

The results provide preliminary support for ADEM as a second-order reflective measure represented by EB, AA, and OP in this sample. Reliability and convergent validity were strong, whereas discriminant validity was qualified by elevated HTMT ratios. AD, EA, RI, and IR were positively associated with ADEM, and IR partially mediated the associations of AD, EA, and RI with ADEM. The instrument should therefore be regarded as promising but not fully validated until the identified overlap is addressed in independent samples.

Acknowledgements:	The authors sincerely appreciate all respondents who voluntarily participated in this study and contributed their time and perspectives. The authors also gratefully acknowledge the Institute for Poverty Research and Management (InsPeK), Universiti Malaysia Kelantan, for supporting and facilitating this research.
Funding Statement:	This research received no specific grant from any funding agency in the public, commercial, or not-for-profit sectors.
Conflict of Interest Statement:	The authors declare that there is no conflict of interest regarding the publication of this paper. All authors approved the final manuscript for submission to the International Journal of Modern Education (IJMOE).
Ethics Statement:	This study was conducted in accordance with institutional ethical research standards. The protocol was approved by Universiti Malaysia Kelantan (Approval ID: UMK.A16.800-6/1/1 (12)). Informed consent was obtained from all participants. Participation was voluntary, and responses were anonymous and used solely for academic purposes.
Author Contribution Statement:	All authors contributed to the conceptualisation, methodology, investigation, interpretation, drafting, and critical revision of the manuscript. All authors read and approved the final version for submission.

References

- Ab Hamid, M. R., Sami, W., & Sidik, M. H. (2017). Discriminant validity assessment: Use of Fornell & Larcker criterion versus HTMT criterion. *Journal of Physics: Conference Series*, 890(1), 012163. <https://doi.org/10.1088/1742-6596/890/1/012163>
- Akintola, A. S., Akintayo, M., Kadri, T., Oforgu, C. M., Michael, M., & Nwanna, M. (2024). Adaptive AI systems in education: Real-time personalised learning pathways for skill development. *Journal of Artificial Intelligence Machine Learning and Data Science*, 3(1), 2489–2494. <https://doi.org/10.51219/JAIMLD/Akinyemi-Sadeeq-Akintola/534>
- Akter, S., Fosso Wamba, S., & Dewan, S. (2017). Why PLS-SEM is suitable for complex modeling? An empirical illustration in Big Data Analytics Quality. *Production Planning & Control*, 28(11–12), 1011–1021. <https://doi.org/10.1080/09537287.2016.1267411>
- Aljuaid, H. (2024). The impact of artificial intelligence tools on academic writing instruction in higher education: A systematic review. *Arab World English Journal (AWEJ) Special Issue on ChatGPT*, 26–55. <https://dx.doi.org/10.24093/awej/ChatGPT.2>
- Aslan, S. (2023). The effect of social comparison on interpersonal problems through self-esteem and emotion dysregulation: Testing a psychological mechanism from an evolutionary framework [Doctoral dissertation, Middle East Technical University]. YÖK National Thesis Center.
- Association of College and Research Libraries. (2016). Framework for information literacy for higher education. American Library Association. https://www.ala.org/sites/default/files/acrl/content/issues/infolit/Framework_ILHE.pdf
- Bandura, A. (1997). Self-efficacy: The exercise of control. W. H. Freeman and Company.
- Banjanovic, E., & Osborne, J. W. (2016). Confidence intervals for effect sizes: Applying bootstrap resampling. *Practical Assessment, Research & Evaluation*, 21(5), 1–20.
- 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>
- Becker, J.-M., Rai, A., Ringle, C. M., & Völckner, F. (2013). Discovering unobserved heterogeneity in structural equation models to avert validity threats. *MIS Quarterly*, 37(3), 665–694. <https://doi.org/10.25300/MISQ/2013/37.3.01>
- Birney, L. B., & McNamara, D. M. (2024). Students' self-efficacy and confidence in technological abilities resulting from participation in "The Curriculum and Community Environmental Restoration Science (STEM + Computer Science)." *Journal of Curriculum and Teaching*, 13(1), 24. <https://doi.org/10.5430/jct.v13n1p24>
- Cepeda-Carrion, G. A., Nitzl, C., & Roldán, J. L. (2018). Mediation analyses in partial least squares structural equation modeling: Guidelines and empirical examples. In J. F. Hair, G. T. M. Hult, C. M. Ringle, & M. Sarstedt (Eds.), *Partial least squares structural equation modeling: Basic concepts, methodological issues and applications* (pp. 201–218). Springer. https://doi.org/10.1007/978-3-319-64069-3_8
- Chin, W. W., Marcolin, B. L., & Newsted, P. R. (2003). A partial least squares latent variable modeling approach for measuring interaction effects: Results from a Monte Carlo simulation study and an electronic-mail emotion/adoption study. *Information Systems Research*, 14, 189–217.
- Cook, R. D. (2000). Detection of influential observation in linear regression. *Technometrics*, 42(1), 65–68.
- Cribari-Neto, F., Ferrari, S. L. P., & Oliveira, W. A. S. C. (2005). Numerical evaluation of tests based on different heteroskedasticity-consistent covariance matrix estimators. *Journal*

- of Statistical Computation and Simulation, 75(8), 611–628. <https://doi.org/10.1080/00949650410001729427>
- Diamantopoulos, A., & Riefler, P. (2011). Using formative measures in international marketing models: A cautionary tale using consumer animosity as an example. *Advances in International Marketing*, 10, 11–30.
- Duarte, P., & Amaro, S. (2018). Methods for modelling reflective-formative second order constructs in PLS: An application to online travel shopping. *Journal of Hospitality and Tourism Technology*, 9(3), 295–313. <https://doi.org/10.1108/JHTT-09-2017-0092>
- Ecker, U. K. H., Lewandowsky, S., Cook, J., Schmid, P., Fazio, L. K., Brashier, N., Kendeou, P., Vraga, E. K., & Amazeen, M. A. (2022). The psychological drivers of misinformation belief and its resistance to correction. *Nature Reviews Psychology*.
- Gabay, R. A., Funa, A. A., & Ricafort, J. D. (2025, August). Generative Artificial Intelligence (GenAI) for academic writing in higher education: A scoping review of applications, challenges, and implications. *Research Square*. <https://doi.org/10.21203/rs.3.rs-7440784/v1>
- Giannakos, M., Horn, M., & Cukurova, M. (2024). Learning, design and technology in the age of AI. *Behaviour & Information Technology*, 44(5), 883–887. <https://doi.org/10.1080/0144929X.2024.2394886>
- Gkintoni, E., Antonopoulou, H., Sortwell, A., & Halkiopoulos, C. (2025). Challenging cognitive load theory: The role of educational neuroscience and artificial intelligence in redefining learning efficacy. *Brain Sciences*, 15(2), 203. <https://doi.org/10.3390/brainsci15020203>
- Gligorea, I., Cioca, M., Oancea, R., & Tudorache, P. (2023). Adaptive learning using artificial intelligence in e-learning: A literature review. *Education Sciences*, 13(12), 1216. <https://doi.org/10.3390/educsci13121216>
- Hair, J. F., Hult, G. T. M., Ringle, C. M., & Ray, S. (2021). Evaluation of reflective measurement models. In J. F. Hair, G. T. M. Hult, C. M. Ringle, M. Sarstedt, N. P. Danks, & S. Ray, *Partial least squares structural equation modeling (PLS-SEM) using R* (pp. 67–87). Springer. https://doi.org/10.1007/978-3-030-80519-7_4
- Hair, J. F., & Alamer, A. (2022). Partial least squares structural equation modeling (PLS-SEM) in second language and education research: Guidelines using an applied example. *Research Methods in Applied Linguistics*, 1(4), 100027. <https://doi.org/10.1016/j.rmal.2022.100027>
- Hair, J. F., Jr., Hult, G. T. M., Ringle, C. M., & Sarstedt, M. (2017). *A primer on partial least squares structural equation modeling (PLS-SEM)* (2nd ed.). SAGE Publications.
- Han, Z., Song, G., Zhang, Y., & Li, B. (2025). Trust the machine or trust yourself: How AI usage reshapes employee self-efficacy and willingness to take risks. *Behavioral Sciences*, 15(8), 1046. <https://doi.org/10.3390/bs15081046>
- Hayes, A. F. (2015). An index and test of linear moderated mediation. *Multivariate Behavioral Research*, 50, 1–22.
- Hickey, G. L., Kontopantelis, E., Takkenberg, J. J. M., & Beyersdorf, F. (2018). Statistical primer: Checking model assumptions with regression diagnostics. *Interactive Cardiovascular and Thoracic Surgery*, 28(1), 1–4. <https://doi.org/10.1093/icvts/ivy207>
- Hoft, J. (2021). Anonymity and confidentiality. In J. C. Barnes & D. R. Forde (Eds.), *The encyclopedia of research methods in criminology and criminal justice*. Wiley-Blackwell. <https://doi.org/10.1002/9781119111931.ch41>
- İlhan, M., Güler, N., Taşdelen Teker, G., & Ergenekon, Ö. (2024). The effects of reverse items on psychometric properties and respondents' scale scores according to different item

- reversal strategies. *International Journal of Assessment Tools in Education*, 11(1), 20–38. <https://doi.org/10.21449/ijate.1345549>
- Kurniawan, D., Samsura, A., & Van Deemen, A. (2024). Global cross-sectional survey on higher education institutions: Analysing the assessment of instrument for evaluating academic expertise on teaching, research and community service. *Masyarakat Kebudayaan dan Politik*, 37(3), 288–299. <https://doi.org/10.20473/mkp.V37I32024.288-299>
- Li, R., Ouyang, S., & Lin, J. (2025). Mediating effect of AI attitudes and AI literacy on the relationship between career self-efficacy and job-seeking anxiety. *BMC Psychology*, 13(1), 454. <https://doi.org/10.1186/s40359-025-02757-2>
- Martinez, M. N., & Bartholomew, M. (2017). What does it “mean”? A review of interpreting and calculating different types of means and standard deviations. *Pharmaceutics*, 9(2), 14. <https://doi.org/10.3390/pharmaceutics9020014>
- Martin, A. J., Nejad, H., Colmar, S. H., & Liem, G. A. D. (2013). Adaptability: How students' responses to uncertainty and novelty predict their academic and non-academic outcomes. *Journal of Educational Psychology*, 105(3), 728–746.
- Méndez-Suárez, M. (2021). Marketing mix modeling using PLS-SEM, bootstrapping the model coefficients. *Mathematics*, 9(15), 1832. <https://doi.org/10.3390/math9151832>
- Mokhtar, S. F., & Yusof, Z. M., & Sapiri. (2023). Confidence intervals by bootstrapping approach: A significance review. *Malaysian Journal of Fundamental and Applied Sciences*, 19, 30–42.
- Ohtani, K. (2000). Bootstrapping R^2 and adjusted R^2 in regression analysis. *Economic Modelling*, 17(4), 473–483. [https://doi.org/10.1016/S0264-9993\(99\)00034-6](https://doi.org/10.1016/S0264-9993(99)00034-6)
- Paas, F., Renkl, A., & Sweller, J. (2003). Cognitive load theory and instructional design: Recent developments. *Educational Psychologist*, 38(1), 1–4.
- Paas, F., & Van Merriënboer, J. J. G. (2020). Cognitive-load theory: Methods to manage working memory load in the learning of complex tasks. *Current Directions in Psychological Science*, 29(4), 362–366. <https://doi.org/10.1177/0963721420922183>
- Peterson, R. A., Kim, Y., & Choi, B. (2020). A meta-analysis of construct reliability indices and measurement model fit metrics. *Methodology*, 16(3), 208–223. <https://doi.org/10.5964/meth.2797>
- 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>
- Premeti, A., Lekka, E., Pilafas, G., & Louka, P. (2024). Exploring working memory deficits in academic learning: Strategies for identification and intervention. *Galore International Journal of Health Sciences and Research*, 9(2), 21–29. <https://doi.org/10.52403/gijhsr.20240203>
- Rashid, N. A., Jamil, N. A., & Ling, T. W. (2025). Rethinking artificial intelligence (AI) in qualitative research. *International Journal of Care Scholars*, 8(2), 211–214. <https://doi.org/10.31436/ijcs.v8i2.468>
- Richardson, M., & Abraham, C. (2012). Psychological correlates of university students' academic performance: A systematic review and meta-analysis. *Psychological Bulletin*, 138(2), 353–387. <https://doi.org/10.1037/a0026838>
- Ringle, C. M., Sarstedt, M., Sinkovics, N., & Sinkovics, R. R. (2023). A perspective on using partial least squares structural equation modelling in data articles. *Data in Brief*, 48, 109074. <https://doi.org/10.1016/j.dib.2023.109074>

- Rönkkö, M., & Cho, E. (2020). An updated guideline for assessing discriminant validity. *Organizational Research Methods*, 23(3), 1–15. <https://doi.org/10.1177/1094428120968614>
- Sandoval, W. A., Greene, J. A., & Bråten, I. (2016). Understanding and promoting thinking about knowledge: Origins, issues, and future directions of research on epistemic cognition. *Review of Research in Education*, 40(1), 1–45. <https://doi.org/10.3102/0091732X16669319>
- Schunk, D., & Dibenedetto, M. (2020). Self-efficacy and human motivation. In A. J. Elliot (Ed.), *Advances in Motivation Science* (Vol. 7, pp. 1–31). Elsevier. <https://doi.org/10.1016/bs.adms.2020.10.001>
- Susnjak, T., Ramaswami, G. S., & Mathrani, A. (2022). Learning analytics dashboard: A tool for providing actionable insights to learners. *International Journal of Educational Technology in Higher Education*, 19(1), 12. <https://doi.org/10.1186/s41239-021-00313-7>
- Sweller, J., & Chandler, P. (1994). Why some material is difficult to learn. *Cognition and Instruction*, 12(3), 185–233.
- Thompson, C., Kim, R. S., Aloe, A., & Becker, B. J. (2017). Extracting the variance inflation factor and other multicollinearity diagnostics from typical regression results. *Basic and Applied Social Psychology*, 39(2), 1–10. <https://doi.org/10.1080/01973533.2016.1277529>
- Tofghi, D. (2020). Bootstrap model-based constrained optimization tests of indirect effects. *Frontiers in Psychology*, 10, 2989. <https://doi.org/10.3389/fpsyg.2019.02989>
- Uzoigwe, E. I. E., & Sanga, I. S. (2025). The epistemological foundations of artificial intelligence (AI) and its relevance for contemporary scholarship. *Pinisi Journal of Art, Humanity and Social Studies*, 5(1), 10–25.
- Voorhees, C., Brady, M. K., Calantone, R., & Ramirez, E. (2015). Discriminant validity testing in marketing: An analysis, causes for concern, and proposed remedies. *Journal of the Academy of Marketing Science*, 44(1), 1–16. <https://doi.org/10.1007/s11747-015-0455-4>
- Walsh, T., Levy, N., Bell, G., Elliott, A., Maclaurin, J., Mareels, I. M. Y., & Wood, F. M. (2019). The effective and ethical development of artificial intelligence: An opportunity to improve our wellbeing. Australian Council of Learned Academies. https://acola.org/wp-content/uploads/2019/07/hs4_artificial-intelligence-report.pdf
- Wineburg, S. S. (1991). Historical problem solving: A study of the cognitive processes used in the evaluation of documentary and pictorial evidence. *Journal of Educational Psychology*, 83(1), 73–87. <https://doi.org/10.1037/0022-0663.83.1.73>
- Xiao, Z., Li, T. W., Karahalios, K., & Sundaram, H. (2023). Inform the uninformed: Improving online informed consent reading with an AI-powered chatbot. In *CHI '23: Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems* (Article 112, pp. 1–17). ACM. <https://doi.org/10.1145/3544548.3581252>
- Zubir, N. F. A., Faizal, D.-R., & Masrom, N. R. (2025). The impact of artificial intelligence (AI) in education systems: Evidence from Malaysia. *Research Square*. <https://doi.org/10.21203/rs.3.rs-6987750/v1>

Appendix

Appendix A. Emotional Balance Scale (10 items)

Measures stress management, emotional regulation, and motivation when using AI.

Response format (all scales): 1 = Strongly Disagree ... 6 = Strongly Agree

#	Item (English [Malay])	Likert (1–6)
1	I feel confident and capable when learning with AI-based tools. [Saya berasa yakin dan berkebolehan apabila belajar menggunakan alatan berasaskan AI.]	1 2 3 4 5 6
2	I can effectively manage my stress when encountering difficulties using AI systems. [Saya dapat menguruskan tekanan saya dengan berkesan apabila menghadapi kesukaran menggunakan sistem AI.]	1 2 3 4 5 6
3	I stay calm even when AI tools do not perform as expected or present challenges. [Saya kekal tenang walaupun alatan AI tidak berfungsi seperti yang dijangka atau menimbulkan cabaran.]	1 2 3 4 5 6
4	I feel intrinsically motivated to persevere even when AI-driven tasks are challenging. [Saya termotivasi secara intrinsik untuk terus gigih walaupun tugas berasaskan AI mencabar.]	1 2 3 4 5 6
5	I can regulate my emotions and maintain a positive outlook while interacting with AI in learning. [Saya dapat mengawal emosi dan mengekalkan pandangan positif semasa berinteraksi dengan AI dalam pembelajaran.]	1 2 3 4 5 6
6	I do not experience significant anxiety when using AI platforms for academic work. [Saya tidak mengalami kebimbangan yang signifikan apabila menggunakan platform AI untuk kerja akademik.]	1 2 3 4 5 6
7	I believe AI can support my emotional well-being and reduce cognitive load during learning. [Saya percaya AI boleh menyokong kesejahteraan emosi saya dan mengurangkan beban kognitif semasa pembelajaran.]	1 2 3 4 5 6
8	I am able to control impulsive reactions when AI gives unexpected or unhelpful responses. [Saya mampu mengawal reaksi impulsif apabila AI memberikan respons yang tidak dijangka atau tidak membantu.]	1 2 3 4 5 6
9	I seek positive emotional engagement when integrating AI into my learning routine. [Saya mencari pengalaman emosi yang positif apabila mengintegrasikan AI ke dalam rutin pembelajaran saya.]	1 2 3 4 5 6
10	I maintain emotional equilibrium by trusting in Allah's plan even when technology is overwhelming. [Saya mengekalkan keseimbangan emosi dengan mempercayai perancangan Allah walaupun teknologi terasa membebankan.]	1 2 3 4 5 6

Appendix B. Intellectual Readiness Scale (10 items)

Measures cognitive engagement, critical thinking, and digital literacy with AI.

#	Item (English [Malay])	Likert (1–6)
1	I understand fundamental principles of how AI tools function in the learning environment. [Saya memahami prinsip-prinsip asas bagaimana alatan AI berfungsi dalam persekitaran pembelajaran.]	1 2 3 4 5 6
2	I can critically evaluate the accuracy and relevance of AI-provided information. [Saya dapat menilai secara kritis ketepatan dan kerelevanan maklumat yang diberikan oleh AI.]	1 2 3 4 5 6
3	I know how to use AI tools effectively to enhance learning and problem-solving. [Saya tahu menggunakan alatan AI dengan berkesan untuk meningkatkan pembelajaran dan penyelesaian masalah.]	1 2 3 4 5 6
4	I can identify and challenge biased, inaccurate, or incomplete AI content. [Saya mampu mengenal pasti dan mencabar kandungan AI yang bias, tidak tepat atau tidak lengkap.]	1 2 3 4 5 6
5	I am confident solving complex academic problems by strategically integrating AI tools. [Saya yakin menyelesaikan masalah akademik yang kompleks dengan mengintegrasikan alatan AI secara strategik.]	1 2 3 4 5 6
6	I actively seek innovative ways to use AI to improve academic performance and understanding. [Saya secara aktif mencari cara baharu menggunakan AI untuk meningkatkan prestasi akademik dan pemahaman.]	1 2 3 4 5 6
7	I understand the limitations of AI in learning and its assistive role. [Saya memahami batasan AI dalam pembelajaran serta peranannya sebagai pembantu.]	1 2 3 4 5 6
8	I engage in metacognition when collaborating with AI on learning tasks. [Saya melibatkan diri dalam metakognisi apabila bekerjasama dengan AI dalam tugas pembelajaran.]	1 2 3 4 5 6
9	I adapt my learning strategies based on AI insights. [Saya bersedia mengubah strategi pembelajaran berdasarkan pandangan AI.]	1 2 3 4 5 6
10	I view AI as a means to expand knowledge, aligned with the Islamic emphasis on seeking 'ilm. [Saya melihat AI sebagai cara untuk mengembangkan ilmu, selaras dengan penekanan Islam terhadap pencarian ilmu.]	1 2 3 4 5 6

Appendix C. Axiological Alignment Scale (10 items)

Measures values-consistent, ethical, and integrity-aligned AI use, including purpose and spiritual well-being.

#	Item (English [Malay])	Likert (1–6)
1	AI-supported learning should align with my moral and Islamic ethical values. [Pembelajaran disokong AI harus selaras dengan nilai moral dan etika Islam saya.]	1 2 3 4 5 6
2	I reflect on the higher purpose (<i>maqasid</i>) of using AI in my education. [Saya merenung tujuan yang lebih tinggi (<i>maqasid</i>) penggunaan AI dalam pendidikan saya.]	1 2 3 4 5 6
3	I avoid dishonest academic practices with AI, upholding <i>amanah</i> . [Saya mengelakkan amalan akademik tidak jujur dengan AI, menjunjung prinsip <i>amanah</i> .]	1 2 3 4 5 6
4	I feel spiritually grounded even when technology is heavily involved in learning. [Saya berasa tenang walaupun teknologi banyak terlibat dalam pembelajaran.]	1 2 3 4 5 6
5	I use AI in ways that respect academic integrity and others' rights, embodying ' <i>adl</i> '. [Saya menggunakan AI dengan menghormati integriti akademik dan hak orang lain, melambangkan keadilan.]	1 2 3 4 5 6
6	Learning with AI should contribute to society and align with divine guidance. [Pembelajaran dengan AI harus menyumbang kepada masyarakat dan sejajar dengan petunjuk ilahi.]	1 2 3 4 5 6
7	I am aware of ethical implications of AI in education (fairness, equity, dignity). [Saya sedar implikasi etika AI dalam pendidikan (keadilan, kesaksamaan, maruah).]	1 2 3 4 5 6
8	I consider AI's impact on <i>tawhid</i> and avoid distractions from spiritual duties. [Saya mempertimbangkan kesan AI terhadap <i>tawhid</i> dan memastikan tidak mengabaikan kewajipan spiritual.]	1 2 3 4 5 6
9	I use AI to promote <i>ihsan</i> in my studies and future contributions. [Saya menggunakan AI untuk mempromosikan <i>ihsan</i> dalam pembelajaran dan sumbangan masa depan.]	1 2 3 4 5 6
10	I seek beneficial knowledge (<i>nafi'</i>) with AI and avoid harmful uses. [Saya berusaha menggunakan AI untuk ilmu yang bermanfaat dan mengelakkan yang memudaratkan.]	1 2 3 4 5 6

Appendix D. AI Adaptability Scale (10 items)

Measures comfort, frequency, and perceived usefulness of AI in learning.

#	Item (English [Malay])	Likert (1–6)
1	I adapt quickly and efficiently to new AI tools in my learning environment. [Saya cepat dan cekap menyesuaikan diri dengan alatan AI baharu.]	1 2 3 4 5 6
2	I feel comfortable using AI as a regular part of daily learning. [Saya selesa menggunakan AI sebagai sebahagian aktiviti pembelajaran harian.]	1 2 3 4 5 6
3	I learn effectively and efficiently with AI-supported platforms. [Saya dapat belajar dengan berkesan dan cekap menggunakan platform disokong AI.]	1 2 3 4 5 6
4	I frequently and willingly use AI for study, research, or assignments. [Saya kerap dan bersedia menggunakan AI untuk belajar, penyelidikan, atau tugas.]	1 2 3 4 5 6
5	I believe AI improves the overall quality and effectiveness of my learning. [Saya percaya AI meningkatkan kualiti dan keberkesanan pembelajaran saya.]	1 2 3 4 5 6
6	I am eager to explore more AI tools and functionalities. [Saya bersemangat meneroka lebih banyak alatan dan fungsi AI.]	1 2 3 4 5 6
7	I am confident I can succeed academically with strategic AI support. [Saya yakin boleh berjaya secara akademik dengan sokongan strategik AI.]	1 2 3 4 5 6
8	I integrate AI flexibly across different learning tasks and contexts. [Saya fleksibel mengintegrasikan AI merentas pelbagai tugas dan konteks pembelajaran.]	1 2 3 4 5 6
9	I proactively seek ways to optimize learning through AI. [Saya secara proaktif mencari cara mengoptimumkan pembelajaran melalui AI.]	1 2 3 4 5 6
10	I regard AI as an indispensable resource for overcoming learning challenges. [Saya menganggap AI sebagai sumber yang sangat diperlukan untuk mengatasi cabaran pembelajaran.]	1 2 3 4 5 6

Appendix E. Epistemological Awareness Scale (10 items)

Measures understanding of knowledge, its validation, and limits when mediated by AI.

#	Item (English [Malay])	Likert (1–6)
1	I critically evaluate AI information rather than accepting it as absolute truth. [Saya menilai secara kritis maklumat AI dan tidak menerimanya sebagai kebenaran mutlak.]	1 2 3 4 5 6
2	I understand AI-generated knowledge is pattern- and data-based with possible biases. [Saya faham pengetahuan dijana AI berasaskan corak dan data serta mungkin mempunyai bias.]	1 2 3 4 5 6
3	I consider the origin, source, and quality of data used to train AI. [Saya mempertimbangkan asal, sumber, dan kualiti data latihan AI.]	1 2 3 4 5 6
4	I recognize AI outputs may be incomplete or not definitively true. [Saya menyedari keluaran AI mungkin tidak menyeluruh atau tidak benar secara muktamad.]	1 2 3 4 5 6
5	I differentiate between knowledge I derive critically and knowledge presented by AI. [Saya membezakan antara pengetahuan hasil pemikiran kritis saya dan yang disampaikan oleh AI.]	1 2 3 4 5 6
6	I reflect on how AI interaction shapes what counts as valid knowledge in my field. [Saya merenung bagaimana interaksi dengan AI membentuk takrif “pengetahuan” yang sah dalam bidang saya.]	1 2 3 4 5 6
7	I am aware AI can “hallucinate” or produce factually incorrect information. [Saya sedar AI boleh “berhalusinasi” atau menghasilkan maklumat yang tidak tepat.]	1 2 3 4 5 6
8	I frequently cross-reference AI answers with textbooks and peer-reviewed articles. [Saya kerap menyemak silang jawapan AI dengan buku teks dan artikel semakan rakan sebaya.]	1 2 3 4 5 6
9	I seek to understand the reasoning or logic behind AI outputs. [Saya berusaha memahami penaaakulan atau logik di sebalik keluaran AI.]	1 2 3 4 5 6
10	I believe human discernment remains paramount when using AI for knowledge. [Saya percaya pertimbangan manusia kekal paling penting ketika menggunakan AI untuk ilmu.]	1 2 3 4 5 6

Appendix F. Ontological Perspective Scale (10 items)

Measures learner identity, authenticity, personal agency, and views of the human–AI relationship in AI-integrated learning.

#	Item (English [Malay])	Likert (1–6)
1	My learning journey remains authentic and meaningful even with substantial AI assistance. [Perjalanan pembelajaran saya kekal asli dan bermakna walaupun dengan bantuan AI yang besar.]	1 2 3 4 5 6
2	AI augments human cognition rather than replacing intrinsic intellectual development. [AI meningkatkan kognisi manusia dan tidak menggantikan pembangunan intelektual intrinsik.]	1 2 3 4 5 6
3	I retain strong personal agency and control over learning, even when relying on AI. [Saya kekal mempunyai agensi dan kawalan peribadi yang kuat terhadap pembelajaran walaupun bergantung pada AI.]	1 2 3 4 5 6
4	I view AI as non-sentient and do not attribute human-like understanding to it. [Saya menganggap AI tidak berperasaan dan tidak menisbahkan pemahaman seperti manusia kepadanya.]	1 2 3 4 5 6
5	I consider long-term implications of AI use for independence and creativity. [Saya mempertimbangkan implikasi jangka panjang penggunaan AI terhadap kebebasan dan kreativiti.]	1 2 3 4 5 6
6	I feel connected to the human aspects of learning (instructors, peers) even with AI mediation. [Saya berasa terhubung dengan aspek kemanusiaan pembelajaran (pengajar, rakan) walaupun dimediasi AI.]	1 2 3 4 5 6
7	I am comfortable with AI shaping practical realities and challenges in learning. [Saya selesa dengan AI yang membentuk realiti praktikal dan cabaran pembelajaran.]	1 2 3 4 5 6
8	My internal drive for learning remains strong despite AI’s growing presence. [Dorongan dalaman saya untuk belajar kekal kuat walaupun kehadiran AI meningkat.]	1 2 3 4 5 6
9	I recognize AI can simulate understanding but is not a conscious being. [Saya menyedari AI boleh mensimulasikan pemahaman tetapi bukan makhluk yang sedar.]	1 2 3 4 5 6
10	My learner identity is enhanced by strategic AI use rather than diminished by over-reliance. [Identiti saya sebagai pelajar dipertingkatkan melalui penggunaan AI secara strategik, bukan dilemahkan oleh ketergantungan berlebihan.]	1 2 3 4 5 6

Appendix G. Resource Integration Awareness Scale (10 items)

Measures how students combine AI tools with conventional academic resources.

#	Item (English [Malay])	Likert (1–6)
1	I actively cross-check AI answers with my textbook or lecture notes. [Saya secara aktif menyemak silang jawapan AI dengan buku teks atau nota kuliah.]	1 2 3 4 5 6
2	I consult traditional scholarly sources before fully accepting AI responses. [Saya merujuk sumber ilmiah tradisional sebelum menerima sepenuhnya respons AI.]	1 2 3 4 5 6
3	I use AI strategically to deepen understanding of complex textbook concepts. [Saya menggunakan AI secara strategik untuk mendalami konsep buku teks yang kompleks.]	1 2 3 4 5 6
4	I compare AI responses with trusted academic references to ensure accuracy. [Saya membandingkan respons AI dengan rujukan akademik yang dipercayai untuk memastikan ketepatan.]	1 2 3 4 5 6
5	I still regard textbooks and lecture notes as my primary, reliable sources. [Saya masih menganggap buku teks dan nota kuliah sebagai sumber utama yang boleh dipercayai.]	1 2 3 4 5 6
6	AI helps me understand difficult readings and challenging concepts. [AI membantu saya memahami bacaan sukar dan konsep mencabar.]	1 2 3 4 5 6
7	I do not rely solely on AI; I maintain a diverse set of learning resources. [Saya tidak bergantung semata-mata pada AI; saya mengekalkan kepelbagaian sumber pembelajaran.]	1 2 3 4 5 6
8	I synthesize information from AI and traditional sources for a comprehensive view. [Saya menggabungkan maklumat daripada AI dan sumber tradisional untuk mendapatkan gambaran yang menyeluruh.]	1 2 3 4 5 6
9	I know when to use AI for quick retrieval versus engaging with primary sources. [Saya tahu bila menggunakan AI untuk capaian pantas berbanding menelaah sumber utama.]	1 2 3 4 5 6
10	I use AI to discover and access relevant traditional academic resources efficiently. [Saya menggunakan AI untuk menemui dan mengakses sumber akademik tradisional dengan cekap.]	1 2 3 4 5 6