



Relationships among AI Literacy, Critical Evaluation of AI-Generated Information, and Academic Integrity among Higher Education Students: A Cross-Sectional Study with Implications for Assessment Practice

Tariq bin Juma bin Mohammad Al-Rasbi¹, Dr. Rabi' bin Al-Mur Al-Dhahli², Prof. Dr. Ahmed Mahmoud Radwan³

¹PhD candidate in Educational Leadership, University of Nizwa, Email: 2097039@uofn.edu.om

²Associate Professor of Educational Administration, University of Nizwa, Email: rabeealthuhli@unizwa.edu.om,

³Professor of Education and Humanities, University of Nizwa, Email: ahmad.rathwan@unizwa.edu.om

Received: 02/06/2026 Revised: 05/07/2026 Acceptance: 13/07/2026 Published: 22/07/2026

ABSTRACT

The integration of generative AI in higher education has sparked learning opportunities as well as concerns over students' ability to assess information generated by AI and ensure academic honesty. This study explored the relationships between the three measures, AI literacy, critical evaluation, and academic integrity, by analyzing the data from a quantitative, cross-sectional, openly available dataset of 972 students enrolled in higher education and 60 seven-point Likert items. The four domains of affective, behavioral, cognitive, and ethical dimensions of AI literacy were measured; critical evaluation was measured by the items of justification and verification; and academic integrity was reverse-scored as academic dishonesty. Descriptive statistics, Cronbach's alpha, Pearson and Spearman correlations, Welch's tests, partial correlations and sensitivity analysis were performed. The relationships were strong and positive between AI literacy and critical evaluation ($r=.749$, $p<.001$) and weak but positive between AI literacy and academic integrity ($r=.171$, $p<.001$). The relationship of ethical AI literacy towards academic dishonesty was the strongest negative one, and institutional difference was also found to be significant. The results suggest a stronger link between AI literacy with the critical evaluation than with integrity-related behaviour. The incorporation of verification, ethical considerations, and disclosure of the use of AI tools into the curriculum and policy of universities is therefore necessary.

Keywords – AI literacy; critical evaluation; academic integrity; generative AI; higher education



1. Introduction

Artificial intelligence (AI), particularly generative AI, has rapidly become embedded in higher education. Students today find information, they create explanations, they summarize academic texts, they get help with their writing and they solve complex problems with the help of AI tools. These advancements have brought about fresh possibilities for individualised learning, access and academic productivity, alongside worries with misinformation, over-reliance, authorship, and ethical application. The spread of AI systems is now more likely than ever before to shape the processes by which students are searching for knowledge, interpreting it, and creating it, and higher education institutions are now looking at students both as people who are able to operate these tools, but also as people who have the knowledge, judgement and ethical awareness needed to use them responsibly. AI literacy, however, is not limited to technical skills but also involves understanding how AI works, what its limitations are and how to evaluate its results as well as considering the social and ethical implications of the use of AI. This multi-dimensional perspective is captured in the Meta AI Literacy Scale (MetAILS) developed by Carolus et al. (2023), which combines the various competencies of AI with psychological and meta-competencies that are relevant to informed and adaptive interactions with AI.

As generative AI spreads, so has the discussion around students' capacities to critically analyse the information they create with AI. AI-generated texts can sound convincing, seem coherent and authoritative even if they are not accurate, biased, incomplete, or fabricated. When students don't have the skills to evaluate, they might take such outputs without adequate verification, particularly if they are in a hurry or are not knowledgeable about the subject. The capacity to evaluate sources and evidence has been found to be one of the key skills in identifying unreliable information, and a critical evaluation has been identified as a separate and required skill in digitally mediated learning environment (Jones-Jang et al., 2021). As part of this competency, in the context of AI, fact checking, comparing output with credible sources, identifying possible bias, recognising uncertainty and deciding when human expertise is required.

Meanwhile, the integration of generative AI has thrown into confusion traditional concepts of academic honesty. Issues discussed are plagiarism from AI, inappropriate outsourcing of the writing to AI, fabrication of references, and turning AI-generated content in as original work. Cotton et al. (2024) believe that institutions need to take action on the above issues by providing more defined guidance, redesigning the assessment process and adopting more educational strategies. This fits into Dawson's (2021) overall notion of the security of assessment should go hand-in-hand with the secure design of valid learning and academic integrity support. Academic integrity extends beyond just conforming to institutional rules—it is about the values of honesty, trust, fairness, respect, responsibility, and courage and is applicable to the use of AI in the academic process (International Center for Academic Integrity, 2021).

While the importance of institutions cannot be denied, the correlation between the students' AI literacy, the capability to critically analyze AI-generated information and the academic integrity is still not fully understood, and this is very important. Students can have a high degree of confidence in the use of AI without fully appreciating its potential drawbacks, or identify



incorrect information and still engage with AI in a manner that is inconsistent with institutional norms. Students can either be able to trust AI without a proper understanding of its limitations or even notice incorrect information yet still use AI in ways that aren't aligned with the institution's norms. On the other hand, better AI literacy can help to foster critical thinking and more responsible academic conduct. From a student's point of view, generative AI seems to be helpful and challenging features, with regard to learning support, accuracy, fairness and ethical boundaries (Chan & Hu, 2023). But all these dimensions are often studied individually, and there is little empirical evidence of the interaction of these dimensions within the same body of students.

In this study the focus is on higher education students and the relations studied using a cross sectional research design. It is only a snapshot of the associations and cannot allow drawing causal conclusions. Additionally, the study is based on self-reported responses that could be subject to social desirability, differing prior experiences with AI, knowledge of institutional policies, and/or different definitions of academic integrity. In addition, the results might be limited to the participating institutions, disciplines or groups of students. However, the study can offer a much-needed cross-sectional snapshot of the current strengths and weaknesses among students in assessing their competencies, practices, and orientations on integrity in a dynamic educational context.

The relevance of the study is that it could contribute to the practice of assessment, curriculum development and institutional policies. It is important to have a holistic perspective on literacy, ethics, teaching, assessment and governance in the context of AI education, rather than teaching the technology alone (Chan, 2023). Current evidence of relationships between AI literacy and critical evaluation and academic integrity could support educator assessments that account for the need to verify, reflect, provide source justification, engage in independent thinking and disclosure of AI use. This evidence can also contribute to designing specific AI literacy programmes that have shown the potential to boost student empowerment, conceptual understanding and ethical awareness (Kong et al., 2023). The findings of this study can help institutions transition from a reactive to a proactive approach to the responsible use of AI and foster meaningful learning and credible evaluation grounded in education.

The research objectives are:

- To examine the level of AI literacy among higher education students.
- To assess students' ability to critically evaluate AI-generated information and their orientation toward academic integrity.
- To determine the relationships among AI literacy, critical evaluation of AI-generated information, and academic integrity.

2. Literature Review

As students are more often facing AI systems in learning, research, and assessment, the concept of AI literacy has become a key issue in higher education. Early conceptual research conceptualizes AI literacy as a set of skills that allows people to understand, use, assess and reflect on AI, but not to simply use an AI-generation tool. Long and Magerko (2020) suggested a



general competency based model that encompasses perception, knowledge of its strengths and weaknesses, interpretation of its products, and ethical and social concerns. Inspired by this, Ng et al. (2021) came up with four interconnected dimensions of AI literacy: AI knowledge and understanding; AI usage and application; AI evaluation and creation; and AI ethics. These frameworks suggest that literacy with AI is multi-dimensional, consisting of technical knowledge, critical assessment, application and ethical understanding.

AI literacy in the field of higher and adult education is still a nascent and poorly understood discipline. In a scoping review, Laupichler et al. (2022) discovered a trend toward increased interest in the abilities related to AI but also noted some inconsistencies in definitions, teaching methods, and evaluation techniques. This conceptual confusion has led to a greater development of standardized measurement instruments. To evaluate the practical understanding, critical appraisal, problem solving, and ethical awareness of people lacking technical backgrounds in AI, the Scale for the Assessment of Non-Experts' AI Literacy was developed (Laupichler et al., 2023). The available instruments, however, are far from uniform regarding their theoretical underpinnings, the populations they are designed for, their dimensions, and their psychometric quality, indicating the need for more instruments to be validated in educational and cultural contexts (Lintner, 2024). The results of a comparative transnational survey also show that the AI literacy of university students varies by demographic, educational, and national context, suggesting that these factors can influence access to and confidence in the use of AI, and that institutional context can impact both. (Mansoor et al., 2024)

As generative AI has been adopted quickly, critical thinking is becoming a necessary part of AI literacy. Generative systems can create fluent and convincing content, but also can create factual inaccuracies, unsupported claims, and fake sources. In a recent study, Walters and Wilder (2023) showed that ChatGPT-generated bibliographic citations were often incorrect and had been completely made up, pointing to the dangers of believing the AI-generated output without proofreading. Lateral reading is a strategy that is relevant to this problem because it involves the user engaging in the process of leaving the original source, accessing independent evidence, and making a judgment based on comparing claims across credible sources (Wineburg & McGrew, 2019). When applied to AI-generated information, these practices can support students in assessing the accuracy of information, the authority behind it, bias, and the support presented in evidentiary sources, all of which helps students evaluate information beyond simple fluency.

Another significant worry for generative AI is its impact on academic honesty and student learning. Perkins (2023) discussed how Large Language Models challenge traditional understandings of writing and authorship and how integrity is traditionally taught and evaluated, and that institutions should rethink their expectations and standards for acceptable forms of support. Some of the same issues are being raised regarding overreliance on AI, the lack of meaningful interaction with disciplinary thinking, cheating and plagiarism, and plagiarism occurring without attribution (Sullivan et al., 2023). The institutional reactions have been mixed at best. A review of the policies of major universities revealed that there was a significant lack of consensus across institutions regarding acceptable use, transparency, assessment design and staff



responsibility, meaning that many policies are still reactive and fragmented (Moorhouse et al., 2023).

International guidelines highlight the need for responses to generative AI to be innovative, while placing a strong focus on the importance of human agency, inclusion, privacy, transparency and ethical governance. To address the potential risks of AI, UNESCO advocates for the creation of clear policies, enhancement of AI skills in schools, safeguarding data security, and the integration of AI to complement, not supplant, human learning and decision-making (Miao & Holmes, 2023). In the context of assessment, this can be done by implementing structured frameworks that outline varying degrees of acceptable AI usage. One such tool that provides such an approach is the Artificial Intelligence Assessment Scale, which sets AI usage levels in assessment activities according to a spectrum from no AI use to purposeful and critical collaboration with AI tools (Perkins et al., 2024). The body of knowledge as a whole indicates that AI literacy, critical evaluation, and academic integrity are closely related concepts. Students need to be familiar with the working of AI and also possess the ability to critically evaluate and use ethical standards in the assessment process. However, there are few studies that have studied all these relationships together, which is why there is a need for empirical research among higher education students.

3. Methodology

3.1 Research Design

The design of this study was empirical, quantitative, observational and cross sectional which aimed to examine the relationship between the variables of AI literacy, critical evaluation of AI-generated information and academic integrity in Higher Education. It was an appropriate design as the data included a measure of the condition at the same time as the sample rather than an intervention, longitudinal sequence, or randomly assigned exposure. Thus, analysis only involved estimating associations and group differences, but no inferences about instructional effects or causal pathways were made. The methodological framework combined educational technology, assessment, epistemic cognition and academic-integrity aspects, which was aligned with the focus of the journal on educational research. Data analysis was done at the individual student respondent level.

3.2 Dataset and Participants

An openly available secondary dataset, AI Literacy, Epistemic Beliefs, and Academic Dishonesty Across Institutional Contexts, released by Mendeley Data, in the form of an Excel spreadsheet and a questionnaire (Yang et al., 2026), was used. The workbook had 972 respondent records, 65 columns and 60 substantive Likert-type items. There were no cases that were excluded due to missing numeric responses or out-of-range values for items, and all records were kept for analysis, meeting the scale-completeness criteria. The sample included 726 respondents who were coded with 0 (institution) and 246 coded as 1 (institution). Demographic variables available were age, gender, major and institution of study type. In addition to not having labels for the coded categories and describing the original recruitment procedure in the supplied files, the sample was



deemed to be a non-probability higher education sample, and the institutional codes were reported without relabeling. The age group of 18–19 years old was the most prevalent group in most respondents. The dataset included higher education students who used AI tools, forming the target population. Data acquisition was made of downloading V1 of the archived data and questionnaire, extracting both files and checking for consistency in their structure before analysis.

3.3 Measurement and Variable Construction

The substantive items were measured with seven-point response scales. AI literacy was measured at the respondent level as the average of 34 items that covered affective, behavioural, cognitive and ethical aspects. The affective part encompassed confidence and attitudes towards AI, the behavioral part practical engagement and usage, the cognitive part knowledge and evaluative understanding, and the ethical part responsible and norm-sensitive usage. Critical evaluation of AI-generated information was measured as the average of three justification items that measure logical appraisal, comparison with alternative sources and verification before accepting. This construct was understood as 'epistemic justification' rather than as an overall assessment of critical thinking. Academic dishonesty was measured as the average of 14 items related to plagiarism, falsification, duplication and unauthorized assistance. All analyses done with reference to academic integrity were reverse-scored (8 minus the dishonesty mean) to make higher scores signify higher integrity. Subscale means were calculated to assess the patterns within each domain. A complete response to all items was essential to be able to score the composite, and all retained records did meet this requirement.

3.4 Data Screening and Scale Evaluation

For data preprocessing, the Python 3 environment was used, with the `artifact_tool` library used for extracting and inspecting workbooks and the NumPy library for handling data and SciPy library for inferential statistics and reproducible statistical computation. Record counts, uniqueness of identifiers, missingness, response ranges, exact duplicate 60-item vectors, within-person variability in response, and extreme response patterns were investigated in screening. Primary analyses used the full sample to maintain transparency and not deleted any data. Potentially low-quality responses were identified when a respondent responded in a single response type (substantive items) throughout the instrument. Exact duplicate response vectors were detected because multiple full-pattern responses can be either due to similarity, copying, batch entry or synthetic generation. For each overall scale and subscale, a Cronbach's alpha was used to estimate internal consistency. The alpha coefficients were interpreted in conjunction with the content of the items and the length of the scale as very high values might reflect redundancy and not better construct validity.

3.5 Statistical Analysis

Coded demographic variables were reported in terms of frequencies and all scale scores were reported in terms of means, standard deviations, minima and maxima. Pearson product-moment correlations were used to measure linear relationships between overall AI literacy, critical



evaluation, academic dishonesty and reverse-scored academic integrity. Spearman rank correlations were used as distribution-invariant tests for the aggregated ordinal responses. One-way independent-samples t tests were used to assess differences between institutions due to the small sample sizes in each and the possibility that the groups might not have been homoscedastic. The Cohen's d was used to measure standardized mean differences. Two-sided tests at the .05 level of significance were also used to determine statistical significance; however, p values were not used exclusively for interpretation but were interpreted in terms of effect size, confidence level, and educational relevance. AI literacy was the central explanatory construct, critical evaluation as an epistemic outcome, and academic integrity as a behavioral-ethical outcome. The ordering directive did not necessarily mean that it was to be interpreted as a temporal or causal ordering.

3.6 Adjusted and Robustness Analyses

The partial correlation and the ordinary least squares regressions involving the focus variables individually and after controlling for institution code, age code, gender code and major code were used to determine the persistence of focal associations. Regression diagnostics were performed, including residual plots, variance inflation plots and influential observation plots, when appropriate. Reliability was assessed by comparing Pearson and Spearman estimates, models without the exact duplicate response vectors, and the models without the respondents who have very little variation in their responses, with the complete sample results, and by conducting sensitivity analyses. The key criterion for validation was consistency with respect to direction, magnitude and inferential conclusions in the specifications. Ethical protection was applied for the secondary nature of the study: the data set included anonymized participant identifiers but not actual identities and analysis was restricted to variables needed for the objectives. Since the documentation provided was not able to provide evidence of original consent or institutional-review procedures, no unsupported claim was made about them. Data was collected on an aggregate basis, analysed locally and reported without re-identification.

4. Results

4.1 Sample and Data-Quality Profile

Final analytic sample consisted of 972 higher education students and 65 variables, of which were 60 substantive questionnaire items on 7-point response scales. The number of respondents in institution code 0 was 726 (74.7%) while the number of respondents in institution code 1 was 246 (25.3%). Most participants were aged 18 or 19 years ($n=933$, 96.0%). Gender codes 1 and 2 represented 574 (59.1%) and 398 (40.9%) respondents, respectively, while major codes 1 and 2 represented 472 (48.6%) and 500 (51.4%) respondents. As the codes themselves did not include the substantive meaning of the variables, all variables coded were left in numerical form.

Table 1 shows that the data set had no missing numeric responses, no values beyond the 1–7 range and no duplicate participant identifiers.

Table 1. Sample characteristics and data-quality indicators



Indicator	Frequency	Percentage
Total respondents	972	100.0
Institution code 0	726	74.7
Institution code 1	246	25.3
Respondents aged 18–19	933	96.0
Complete numeric records	972	100.0
Out-of-range item values	0	0.0
Duplicate participant identifiers	0	0.0
Duplicate full-response rows beyond first occurrence	67	6.9
Respondents using no more than two response categories	155	15.9
All AI-literacy items scored 7	87	9.0
All dishonesty items scored 1	125	12.9
Combined AI-literacy=7 and dishonesty=1 pattern	34	3.5

The screening results showed 67 records with the same response vector as one or more of the first 60 records. There were 25 identical response patterns, the largest of which occurred. Furthermore, 155 respondents responded to no more than two categories on each substantive item – giving limited within-person variability of response. These observations were kept in the main analysis and were assessed in the robustness analysis.

The institutional composition of the sample is displayed in Figure 1.

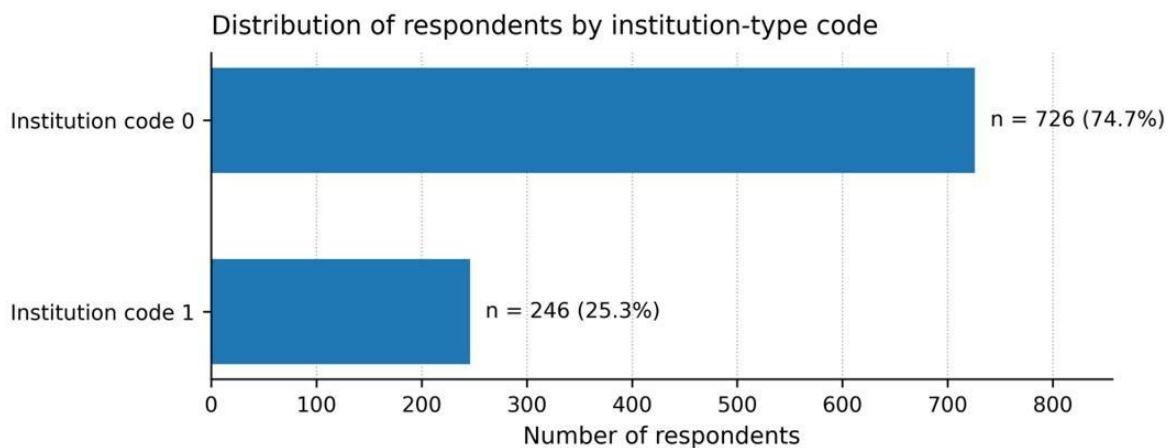


Figure 1. Distribution of respondents by institution code

4.2 Descriptive Statistics

Overall AI literacy had a mean of 5.345 (SD = 1.075, 95% CI [5.277, 5.413]). Among its dimensions, ethical AI literacy produced the highest mean (M = 5.677, SD = 1.167), followed by affective AI literacy (M = 5.359, SD = 1.135), behavioral AI literacy (M = 5.226, SD = 1.170), and cognitive AI literacy (M = 5.127, SD = 1.173). Critical evaluation of AI-generated information had a mean of 5.498 (SD = 1.131, 95% CI [5.427, 5.569]).

The mean academic dishonesty score was 3.063 (SD = 1.519, 95% CI [2.967, 3.158]). Reverse



scoring yielded an academic integrity mean of 4.937 (SD = 1.519, 95% CI [4.842, 5.033]). Unauthorized assistance had the highest mean among the dishonesty domains (M = 3.345, SD = 1.609), followed by falsification (M = 3.274, SD = 1.678), duplication (M = 2.928, SD = 1.655), and plagiarism (M = 2.847, SD = 1.685).

Descriptive statistics and reliability coefficients of the study constructs are presented in Table 2.

Table 2. Descriptive statistics and internal consistency of the study measures

Construct	Items	Mean	SD	95% CI	Cronbach's α
Overall AI literacy	34	5.345	1.075	[5.277, 5.413]	0.986
Affective AI literacy	12	5.359	1.135	[5.287, 5.430]	0.978
Behavioral AI literacy	8	5.226	1.170	[5.152, 5.299]	0.958
Cognitive AI literacy	7	5.127	1.173	[5.053, 5.201]	0.943
Ethical AI literacy	7	5.677	1.167	[5.604, 5.751]	0.965
Critical evaluation	3	5.498	1.131	[5.427, 5.569]	0.914
Academic dishonesty	14	3.063	1.519	[2.967, 3.158]	0.965
Academic integrity	14	4.937	1.519	[4.842, 5.033]	0.965
Plagiarism	5	2.847	1.685	[2.741, 2.953]	0.963
Falsification	3	3.274	1.678	[3.169, 3.380]	0.875
Duplication	3	2.928	1.655	[2.824, 3.033]	0.908
Unauthorized assistance	3	3.345	1.609	[3.244, 3.447]	0.842

All focal composites covered the full observed range from 1 to 7. The means of the principal constructs are displayed in Figure 2.

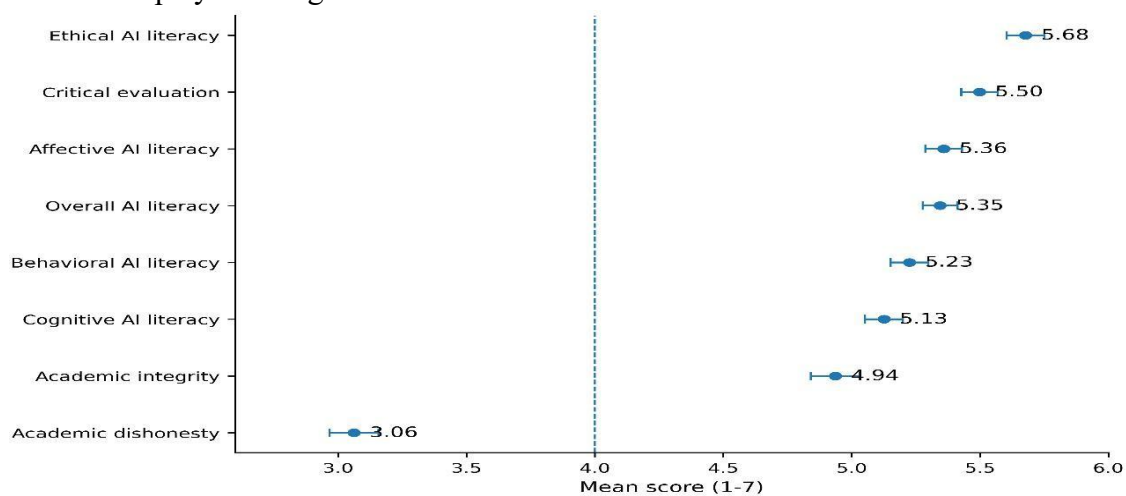


Figure 2. Mean scores of the principal study constructs



4.3 Reliability of the Measures

All internal-consistency coefficients were greater than 0.84 for all scales and subscales. The coefficient ($\alpha = 0.986$) was the highest for overall AI literacy. The coefficients of affective, behavioral, cognitive and ethical dimensions were 0.978, 0.958, 0.943 and 0.965, respectively. The results of the critical evaluation showed that $\alpha = 0.914$, and the overall academic dishonesty measure resulted in $\alpha = 0.965$.

The reliability for the dishonesty domains ranged from 0.842 (for unauthorized assistance) to 0.963 (for plagiarism). The coefficients of falsification and duplication were 0.875 and 0.908, respectively. The coefficients were computed from the full sample of 972 respondents since no item-level data was missing.

4.4 Relationship between AI Literacy and Critical Evaluation

The overall AI literacy was significantly and positively associated with critical evaluation of AI-generated information ($r = 0.749$, 95% CI [0.720, 0.776], $p < 0.001$). The Spearman coefficient was $\rho = 0.717$ ($p < 0.001$). The partial correlation remained significant after adjusting for institution code, age code and gender code and major code ($r = 0.739$, 95% CI [0.709, 0.766], $p < 0.001$).

The analyses conducted at the dimension level revealed the highest correlation between affective AI literacy and critical evaluation, at 0.730 ($p < .001$). Ethical and behavioral AI literacy coefficients were similar ($r = 0.695$, $p < 0.001$; $r = 0.694$, $p < 0.001$, respectively). Positive relationships were also found between cognitive AI literacy and critical evaluation ($r = 0.644$, $p < 0.001$).

The results of the focal correlation estimates are summarized in Table 3.

Table 3. Associations of AI literacy with critical evaluation and academic integrity outcomes

Predictor and outcome	Pearson r	95% CI	Spearman ρ	p
Overall AI literacy–critical evaluation	0.749	[0.720, 0.776]	0.717	<0.001
Affective AI literacy–critical evaluation	0.730	[0.699, 0.758]	0.700	<0.001
Behavioral AI literacy–critical evaluation	0.694	[0.660, 0.725]	0.662	<0.001
Cognitive AI literacy–critical evaluation	0.644	[0.605, 0.679]	0.606	<0.001
Ethical AI literacy–critical evaluation	0.695	[0.661, 0.726]	0.657	<0.001
Overall AI literacy–academic dishonesty	–0.171	[–0.232, –0.109]	–0.258	<0.001
Overall AI literacy–academic integrity	0.171	[0.109, 0.232]	0.258	<0.001
Overall AI literacy–plagiarism	–0.221	[–0.280, –0.161]	–0.302	<0.001



Overall AI literacy– falsification	–0.093	[–0.155, –0.030]	–0.122	0.004
Overall AI literacy–duplication	–0.182	[–0.242, –0.120]	–0.255	<0.001
Overall AI literacy– unauthorized assistance	–0.084	[–0.146, –0.021]	–0.105	0.009

All AI-literacy dimensions were positively related to critical evaluation, with coefficients ranging from 0.644 to 0.749. Their relative magnitudes are displayed in Figure 3.

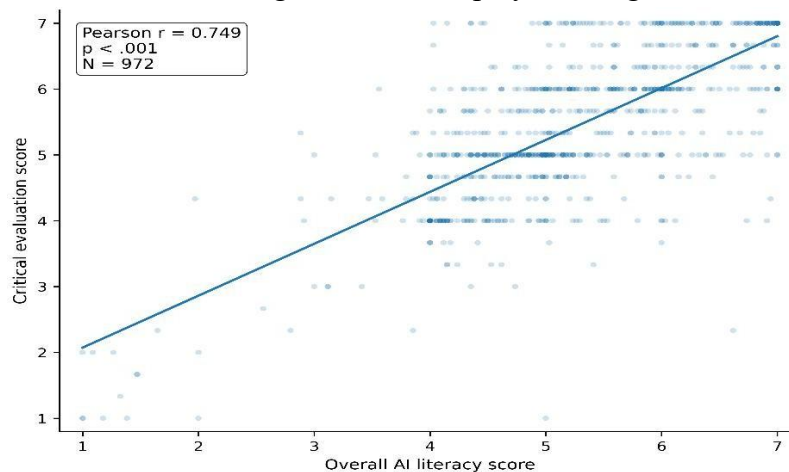


Figure 3. Correlations between AI-literacy measures and critical evaluation

4.5 Relationship between AI Literacy and Academic Integrity

Overall AI literacy was negatively associated with academic dishonesty ($r = -0.171$, 95% CI $[-0.232, -0.109]$, $p < 0.001$). The equivalent association with reverse-scored academic integrity was positive ($r = 0.171$, 95% CI $[0.109, 0.232]$, $p < 0.001$). The rank-order correlation coefficients were respectively $\rho = -0.258$ for dishonesty and $\rho = 0.258$ for integrity ($p < 0.001$ for both).

When institution code, age code, gender code, and major code were held constant, AI literacy continued to have a negative association with academic dishonesty (partial $r = -0.123$, 95% CI $[-0.185, -0.061]$, $p < 0.001$). The corresponding adjusted association with academic integrity was partial $r = 0.123$ (95% CI $[0.061, 0.185]$, $p < 0.001$).

Ethical AI literacy had the largest negative correlation with academic dishonesty at the subdimension level ($r = -0.282$, $p < 0.001$). The affective, behavioral, and cognitive AI literacy had coefficients of -0.153 , -0.122 , and -0.088 , respectively. The cognitive coefficient was still



statistically significant ($p = 0.006$).

There were also differences in dishonesty domain. The absolute relationship that was the largest was between overall AI literacy and plagiarism ($r = -0.221$, $p < 0.001$). The coefficients were found to be -0.182 (duplication), -0.093 (falsification), and -0.084 (unauthorized assistance). All the coefficient values were statistically significant.

4.6 Institutional Comparisons and Robustness

The overall AI literacy level was higher for respondents with institution code 0 compared to those with institution code 1. Mean AI literacy was 5.484 ($SD = 1.096$) for code 0 and 4.936 ($SD = 0.896$) for code 1, producing a difference of -0.548 for code 1 relative to code 0 (95% CI $[-0.686, -0.411]$, Welch $t(512.10) = 7.823$, $p < 0.001$, $d = -0.523$).

Also, there was a higher critical evaluation among those who responded to institution code 0. The means were 5.620 ($SD = 1.158$) and 5.138 ($SD = 0.961$), respectively, with a difference of -0.482 (95% CI $[-0.629, -0.335]$, Welch $t(504.38) = 6.441$, $p < 0.001$, $d = -0.434$).

The prevalence of academic dishonesty was greater in institution code 1. Mean scores were 2.796 ($SD = 1.565$) for code 0 and 3.850 ($SD = 1.029$) for code 1, yielding a difference of 1.054 (95% CI $[0.882, 1.226]$, Welch $t(645.74) = -12.030$, $p < 0.001$, $d = 0.728$). There was a difference in academic integrity that was reverse scored, with a negative score of -1.054 ($d = -0.728$).

Comparisons between the institutions and sensitivity estimates are summarized in Table 4.

Table 4. Institution-code comparisons and robustness estimates

Analysis	Outcome	N	Estimate	95% CI	p
Institution comparison	AI-literacy difference, code 1–0	972	-0.548	$[-0.686, -0.411]$	<0.001
Institution comparison	Critical-evaluation difference, code 1–0	972	-0.482	$[-0.629, -0.335]$	<0.001
Institution comparison	Dishonesty difference, code 1–0	972	1.054	$[0.882, 1.226]$	<0.001
Institution comparison	Integrity difference, code 1–0	972	-1.054	$[-1.226, -0.882]$	<0.001
Full sample	AI literacy–critical evaluation	972	0.749	$[0.720, 0.776]$	<0.001
Duplicate patterns removed	AI literacy–critical evaluation	905	0.735	—	<0.001



Low-variation responses removed	AI literacy–critical evaluation	817	0.706	—	<0.001
Both screening criteria applied	AI literacy–critical evaluation	794	0.704	—	<0.001
Full sample	AI literacy–academic integrity	972	0.171	[0.109, 0.232]	<0.001
Duplicate patterns removed	AI literacy–academic integrity	905	0.121	—	<0.001
Low-variation responses removed	AI literacy–academic integrity	817	0.180	—	<0.001
Both screening criteria applied	AI literacy–academic integrity	794	0.177	—	<0.001

The critical-evaluation coefficient ranged from 0.704 to 0.749 across the screened samples. The academic-integrity coefficient ranged from 0.121 to 0.180 and remained positive under every screening specification.

The stability of the focal estimates across the four analytic samples is illustrated in Figure 4.

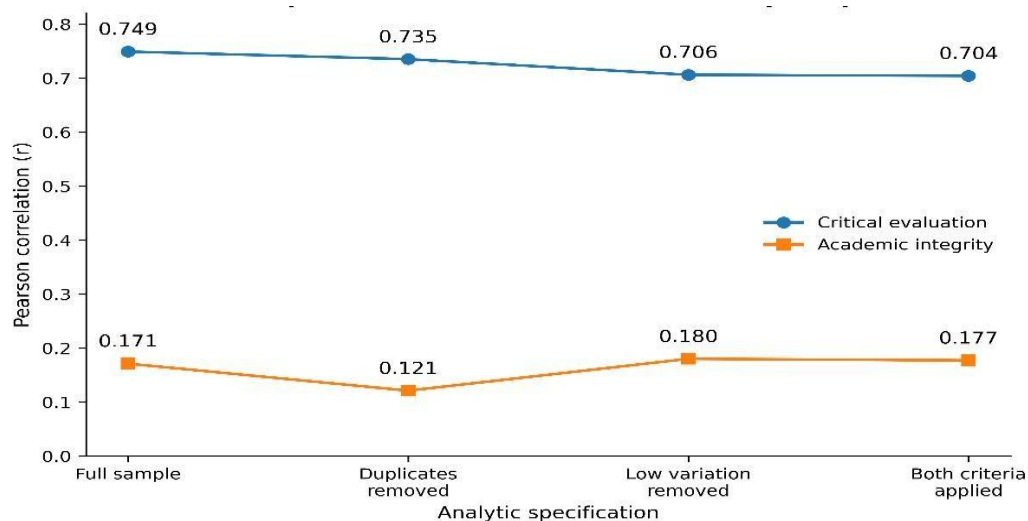


Figure 4. Sensitivity estimates for the focal correlations



5. Discussion

The results suggest that the students' knowledge of AI literacy is closely linked to their critical evaluation competencies related to AI-generated information, while the link with academic integrity is positive but relatively low. The overall AI literacy was found to be highly correlated with critical evaluation ($r = .749$) and this correlation was still significant when controlling for institution, age, gender code and major code (partial $r = .739$). It is stable upon all three analyses-Pearson, Spearman, adjusted and sensitivity analyses- which indicates that it was not sensitive to any one analytical specification. Pupils who felt they knew more about AI, were more confident about their understanding, used it more, and were aware of the ethical issues were also more inclined to check the accuracy of AI-generated information, cross-check AI content with other sources, and evaluate the rationality of AI-generated content. This aligns with Carolus et al. (2023) who suggested a multi-dimensional perspective of AI literacy, including technical skills, reflective skills and meta-competencies. The critical-evaluation score should, however, be interpreted with caution as it reflects perceived ability as opposed to actual performance as reflected in the verification tasks, but it is a relatively high score.

The partial correlation between academic integrity and AI literacy was statistically significant but weak ($r = .171$; adjusted partial $r = .123$). This implies that while literacy can help to promote responsible behaviour, it cannot stop misbehaviour on its own. Assessment pressure, peer norms, clarity of policy, opportunity, and perceptions of fairness are also factors that influence academic integrity. This finding of a stronger negative association between the ethical dimension of AI literacy and academic dishonesty ($r = -.282$), contrasted with the affective, behavioral, and cognitive dimensions is important. The sense of right or wrong, and the knowledge of appropriate behavior seems to be more important in terms of integrity than confidence or technical expertise. The association between AI literacy and plagiarism was the strongest, with relationships with smaller with falsification and unauthorized assistance. It is, therefore, important that teacher's literacy education is linked to authorship, attribution, verification, and disclosure of AI use explicitly. This is especially important when students are using the generative AI technology for legitimate academic work, but then end up using it to generate references that are inaccurate or fabricated, which can result in integrity violations (Walters & Wilder, 2023).

The institutional comparisons also show that competencies and behaviour in the field of AI is context-dependent. The results from the two institutional code respondents differ, with respondents in institution code 0 having higher AI literacy and critical evaluation while higher academic dishonesty is observed in respondents in institution code 1. The size of the effects were moderate for AI literacy and critical evaluation and relatively large for dishonesty. The fact that the differences between these contexts are similar to those found in the literature on students' AI literacy (Mansoor et al., 2024) supports evidence of the variability of AI literacy across educational and national settings. However, the dataset did not offer an explanation of what the institutional codes were or information about differences in curriculum, policy, assessment, or access to AI tools, or student support. The differences observed, then, cannot be explained in terms of any specific institutional characteristics.



The implications of these findings are for curriculum, assessment and policy. AI literacy should not just be about working with AI tools, but should be part of programmes and integrated into all university processes and activities, including the development and verification of sources, understanding bias, and making decisions about AI. Assessment instructions should specify what kind of work is allowed (e.g., AI can be used for generating ideas, but the final product must be the student's own), include clear identification of the use of AI, and ask for drafts or source logs, oral explanations or reflections of the learning process. These practices can help to minimize confusion and continue a meaningful interaction with AI. They also support the idea that institutions should restate and assess their programs to adhere to clear guidance instead of just detection and punishment, primarily (Cotton et al., 2024).

Several limitations qualify the conclusions. Neither causal interpretations are possible because of the cross-sectional design nor self-report measures are free from social desirability, overconfidence and common-method bias. The majority of the participants were concentrated in the age group 18-19, the various institutional groups were not equal and the demographic and institutional codes were not substantively labelled. A single item judgment task was used to assess critical evaluation, instead of the objective task. Data-quality issues were noted in the consistency of response patterns and records with little variation, but were not present in the focal associations when these were removed. Furthermore, the very high reliability coefficients could be due to the redundant nature of the items and not to greater validity.

Future studies should employ longitudinal or experimental design to show if teaching students about AI-literacy positively affects their verification behavior and decreases academic dishonesty. Students should be evaluated on their ability to detect factual errors, quotation or paraphrasing of fabricated sources, bias, and unsupported claims by means of performance assessments. More diverse institutions, disciplines, age groups and national settings should be studied as well. Ethical AI literacy could be explored in mediation models to understand the relationship between wider literacy and integrity, and assessment cultures, policies, and student motivations could be explored in mixed-methods research to clarify the relationship between them and ethical AI use.

6. Conclusion

This study examined relationships among AI literacy, critical evaluation of AI-generated information, and academic integrity among higher education students. A strong correlation was seen between critical evaluation and AI literacy, with students having high levels of knowledge, confidence, engagement, and ethical understanding more likely to verify content, compare sources, and evaluate AI produced content. AI literacy, on the other hand, had a positive, though weak, connection with academic integrity, indicating that advancements in AI skills do not automatically lead to ethical academic conduct. Ethical AI literacy was the strongest indicator of lower academic dishonesty, with institutional differences suggesting that the competencies and practices of integrity in terms of AI are context-dependent. The results suggest that AI literacy should be integrated into the curriculum and into the assessment process at universities, and not taught as a technical skill. Assessment tasks should involve source verification, transparency for the use of AI, provide documentation of the process, require a reflective justification, and evidence of independent



reasoning. It is also important for institutions to offer explicit instructions regarding appropriate vs inappropriate uses of AI in their respective fields. Longitudinal, experimental and mixed-method designs should be adopted in future research to confirm the effectiveness of targeted AI-literacy instruction and its potential for enhancing the ability to verify and reducing misconduct. Objectively moderated critical-evaluation tasks, indicators of behavioral integrity, various institutions, disciplines, and age groups should be included in studies, and the moderating role of ethical AI literacy between AI competence (broad) and academic integrity should be tested.

References

1. Carolus, A., Koch, M. J., Straka, S., Latoschik, M. E., & Wienrich, C. (2023). MAILS—Meta AI literacy scale: Development and testing of an AI literacy questionnaire based on well-founded competency models and psychological change- and meta-competencies. *Computers in Human Behavior: Artificial Humans*, 1(2), Article 100014. <https://doi.org/10.1016/j.chbah.2023.100014>
2. 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, Article 38. <https://doi.org/10.1186/s41239-023-00408-3>
3. Chan, C. K. Y., & Hu, W. (2023). Students' voices on generative AI: Perceptions, benefits, and challenges in higher education. *International Journal of Educational Technology in Higher Education*, 20, Article 43. <https://doi.org/10.1186/s41239-023-00411-8>
4. Cotton, D. R. E., Cotton, P. A., & Shipway, J. R. (2024). Chatting and cheating: Ensuring academic integrity in the era of ChatGPT. *Innovations in Education and Teaching International*, 61(2), 228–239. <https://doi.org/10.1080/14703297.2023.2190148>
5. Dawson, P. (2021). *Defending assessment security in a digital world: Preventing e-cheating and supporting academic integrity in higher education*. Routledge. <https://doi.org/10.4324/9780429324178>
6. International Center for Academic Integrity. (2021). *The fundamental values of academic integrity* (3rd ed.). <https://academicintegrity.org/aws/ICAI/pt/sp/fundamental>
7. Jones-Jang, S. M., Mortensen, T., & Liu, J. (2021). Does media literacy help identification of fake news? Information literacy helps, but other literacies do not. *American Behavioral Scientist*, 65(2), 371–388. <https://doi.org/10.1177/0002764219869406>
8. Kong, S.-C., Cheung, W. M.-Y., & Zhang, G. (2023). Evaluating an artificial intelligence literacy programme for developing university students' conceptual understanding, literacy, empowerment and ethical awareness. *Educational Technology & Society*, 26(1), 16–30. [https://doi.org/10.30191/ETS.202301_26\(1\).0002](https://doi.org/10.30191/ETS.202301_26(1).0002)



9. Laupichler, M. C., Aster, A., Schirch, J., & Raupach, T. (2022). Artificial intelligence literacy in higher and adult education: A scoping literature review. *Computers and Education: Artificial Intelligence*, 3, Article 100101. <https://doi.org/10.1016/j.caeai.2022.100101>
10. Laupichler, M. C., Aster, A., Haverkamp, N., & Raupach, T. (2023). Development of the “Scale for the assessment of non-experts’ AI literacy”—An exploratory factor analysis. *Computers in Human Behavior Reports*, 12, Article 100338. <https://doi.org/10.1016/j.chbr.2023.100338>
11. Lintner, T. (2024). A systematic review of AI literacy scales. *npj Science of Learning*, 9, Article 50. <https://doi.org/10.1038/s41539-024-00264-4>
12. Long, D., & Magerko, B. (2020). What is AI literacy? Competencies and design considerations. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems* (pp. 1–16). Association for Computing Machinery. <https://doi.org/10.1145/3313831.3376727>
13. Mansoor, H. M. H., Bawazir, A., Alsabri, M. A., Alharbi, A., & Okela, A. H. (2024). Artificial intelligence literacy among university students—A comparative transnational survey. *Frontiers in Communication*, 9, Article 1478476. <https://doi.org/10.3389/fcomm.2024.1478476>
14. Miao, F., & Holmes, W. (2023). Guidance for generative AI in education and research. UNESCO. <https://doi.org/10.54675/EWZM9535>
15. Moorhouse, B. L., Yeo, M. A., & Wan, Y. (2023). Generative AI tools and assessment: Guidelines of the world’s top-ranking universities. *Computers and Education Open*, 5, Article 100151. <https://doi.org/10.1016/j.caeo.2023.100151>
16. 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, Article 100041. <https://doi.org/10.1016/j.caeai.2021.100041>
17. Perkins, M. (2023). Academic integrity considerations of AI large language models in the post-pandemic era: ChatGPT and beyond. *Journal of University Teaching & Learning Practice*, 20(2), Article 7. <https://doi.org/10.53761/1.20.02.07>
18. Perkins, M., Furze, L., Roe, J., & MacVaugh, J. (2024). The Artificial Intelligence Assessment Scale (AIAS): A framework for ethical integration of generative AI in educational assessment. *Journal of University Teaching & Learning Practice*, 21(6). <https://doi.org/10.53761/q3azde36>
19. Sullivan, M., Kelly, A., & McLaughlan, P. (2023). ChatGPT in higher education:



Considerations for academic integrity and student learning. *Journal of Applied Learning & Teaching*, 6(1), 31–40. <https://doi.org/10.37074/jalt.2023.6.1.17>

20. Walters, W. H., & Wilder, E. I. (2023). Fabrication and errors in the bibliographic citations generated by ChatGPT. *Scientific Reports*, 13, Article 14045. <https://doi.org/10.1038/s41598-023-41032-5>

21. Wineburg, S., & McGrew, S. (2019). Lateral reading and the nature of expertise: Reading less and learning more when evaluating digital information. *Teachers College Record*, 121(11), 1–40. <https://doi.org/10.1177/016146811912101102>

22. Yang, G., Kim, J., Chen, J., Hao, K., & Li, N. (2026). AI literacy, epistemic beliefs, and academic dishonesty across institutional contexts (Version 1) [Data set]. Mendeley Data. <https://doi.org/10.17632/252r>