

Metacognition and AI Literacy for Generative AI Use in Core Computing Courses: Evaluation and Findings

Rocio Contreras-Aguila, Universidad Andres Bello

Rocío Contreras Águila, M.Sc., is a computer engineer, academic, and consultant specializing in project management and digital transformation. She is currently a lecturer at Universidad Andrés Bello (UNAB), Universidad de Las Américas (UDLA), and Duoc UC, where she teaches software engineering, project management, software quality, databases, and digital transformation. She holds a Master's in Higher Education Teaching, an MBA in Business Administration, and a Master's in Big Data and Business Intelligence. She has extensive experience in the technology sector, leading organizational reengineering and implementing best practices. She is a co-founder of Bitbiosis, focused on software development and business process improvement. Her expertise includes agile and traditional project management, software quality, data analytics, and AI in education. Her research and academic work focus on integrating innovative methodologies in higher education and promoting digital competencies, critical thinking, and metacognitive awareness in AI-mediated learning environments.

Prof. Maria Elena Truyol, Universidad Andres Bello

María Elena Truyol, Ph.D., is full professor and researcher of the Universidad Andrés Bello (UNAB). She graduated as a physics teacher (for middle and high school), physics (M.Sc.), and a Ph.D. in Physics at Universidad Nacional de Córdoba, Argentina. In 2013, she obtained a three-year postdoctoral position at the Universidade de Sao Paulo, Brazil. She focuses on educational research, physics education, problem solving, instructional material design, teacher training, and gender studies. She teaches undergraduate courses in environmental management, energy, and the fundamentals of industrial processes at the School of Engineering, UNAB. She currently coordinates the Educational and Academic Innovation Unit at the School of Engineering (UNAB). She is engaged in continuing teacher training in active learning methodologies at the three campuses of the School of Engineering (Concepción, Viña del Mar, and Santiago, Chile). She authored several manuscripts in the science education area, joined several research projects, participated in international conferences with oral presentations and keynote lectures, and served as a referee for journals, funding institutions, and associations.

Metacognition and AI Literacy for Generative AI Use in Core Computing Courses: Evaluation and Findings

Abstract

The incorporation of artificial intelligence (AI) into higher education not only introduces tools that automate tasks and support problem-solving; but it also reshapes teaching and learning approaches. In engineering—and particularly within the computing curriculum—this transformation is critical because professional preparation demands both technics competencies and critical–reflective skills. Courses such as programming and databases therefore provide a strategic setting: students work with formal languages, logical structures, and complex problems, an ideal context to examine how AI literacy—understood as the ability to understand, apply, and critically evaluate these technologies—relates to metacognitive awareness, defined as the capacity to plan, monitor, and regulate one’s own learning. This articulation is crucial for promoting autonomous learning and fostering an ethical, critical engagement with technology. This study examines that relationship in undergraduate students enrolled in programming and database courses using a two-block questionnaire: the Meta AI Literacy Scale (MAILS), which assesses conceptual knowledge, use/application, ethics, self-efficacy, socioemotional self-competence, and notions of AI creation; and the Metacognitive Awareness in AI Use (MAI-U), an adaptation of classic metacognition frameworks contextualized to AI use and integrating technics components (system limitations and trustworthiness) and ethical components (bias, transparency, and responsibility). Using a cross-sectional design and quantitative approach, the instrument was administered in class, with informed consent, to a sample of 129 students at a private engineering school in Chile. Subscale scores and composite indices (e.g., global MAI) were constructed, reliability and descriptive statistics were estimated, and associations were analyzed via correlations and linear regressions with robust standard errors and collinearity diagnostics, modeling as outcomes AI use and application and AI creation, and including as predictors global MAI, AI self-efficacy, and AI self-competence. Preliminary results indicate high levels of self-efficacy and conceptual AI literacy, as well as comparatively lower levels in creation and detection. The observed associations show a robust convergence between AI metacognition (MAI) and the various literacy facets (knowing, using, detecting, and ethics). Provisional regression models suggest that MAI is the primary predictor of AI use and application as well as conceptual AI knowledge, with self-efficacy providing an additional contribution; perceived competence shows no unique effect once the former are controlled. Taken together, these findings enable concrete curricular actions in computing education: strengthening metacognition applied to AI in early courses to boost use and application; complementing with strategies that reinforce self-efficacy; and addressing gaps in creation and detection through authentic tasks and clear quality criteria. This alignment enables programs to connect learning outcomes, assessments, and instructional support, allowing them to monitor cohort progress and make informed decisions based on evidence. Ultimately, the results support a competent, ethical, and sustainable adoption of AI in the education of future engineers.

Keywords: artificial intelligence, AI literacy, metacognition, ethics, engineering, programming.

Introduction

The emergence of generative artificial intelligence (AI) in higher education has expanded the range of tools available to support learning, especially in disciplines focused on problem-solving [1]. In engineering and computing, tools such as conversational assistants and code generators can improve productivity, support guided learning, and reduce initial friction in complex tasks [2]. However, their use can also produce unintended effects, such as technological dependence, confirmation biases, uncritical acceptance of plausible but incorrect answers, or displacement of essential cognitive processes like debugging, verification, and argumentation [3], [4].

In this context, AI literacy has been proposed as a set of skills that enable understanding, using, evaluating, and making informed decisions regarding AI systems [5], [6]. This literacy includes dimensions such as conceptual knowledge, contextual use, error and bias detection, ethical awareness, and psychological aspects like self-efficacy and critical self-competence [7], [8].

Meanwhile, metacognition—understood as the ability to plan, monitor, and regulate one's own learning—consistently correlates with deep learning, knowledge transfer, and academic performance [9], [10]. In AI-mediated contexts, metacognition becomes especially critical, as the student must set goals, evaluate the quality and reliability of generated responses, compare results with external criteria, and adjust strategies when the AI fails or induces errors [11], [12].

Despite the growing interest in these topics, empirical evidence that jointly addresses AI literacy and metacognitive awareness in core computing courses remains scarce, especially in Latin American contexts [13]. This study contributes to this gap by exploring the descriptive profile of AI literacy and metacognition, their associations, and the incremental predictive contribution of metacognition, self-efficacy, and self-competence.

Literature Review

The integration of artificial intelligence, particularly generative AI, into higher education has markedly transformed pedagogical and learning methodologies, especially within disciplines dedicated to addressing complex problems, such as engineering and computing [1], [2]. Although these models have facilitated the reduction of preliminary technical obstacles and broadened access to immediate cognitive assistance, they have concurrently begun to supplant traditional cognitive processes integral to problem-solving, including hypothesis formulation, verification of results, and disciplinary argumentation [3], [14].

Beyond its instrumental potential, AI introduces cognitive, ethical, and pedagogical challenges that necessitate robust conceptual frameworks for its integration into education [4], [15]. In the absence of explicit pedagogical mediation, intensive utilization of generative AI may result in phenomena such as cognitive dependence or the deterioration of critical thinking skills [3], [14]. In response to this situation, recent scholarly literature has begun to delineate a set of essential cognitive and socio-

emotional skills aimed at guiding a formative, critical, and self-regulated application of AI within university contexts.

Cognitive and socio-emotional skills for critical literacy in AI

AI literacy consists of multiple competencies that enable understanding, using, and critically assessing AI systems [5]–[7], [13]. This literacy encompasses technical, cognitive, ethical, and socio-emotional dimensions, with a particular emphasis on human agency and accountability in decision-making processes [8], [15]. In fundamental computing disciplines such as programming and database management, these skills hold particular significance, as students are required to assess the quality and validity of solutions produced by AI systems within contexts governed by rigorous formal standards [12], [16].

Within this framework, metacognition emerges as a pivotal ability that has been consistently associated with deep and self-regulated learning [9], [10], [17]. In AI-mediated contexts, it is essential for critically evaluating generated responses, comparing them with external evidence, and adjusting strategies when errors or biases are detected [11], [14]. In fact, recent studies indicate that metacognition is a central predictor of effective use and conceptual knowledge in AI, even when controlling for motivational and sociodemographic variables [16], [18].

Closely related to the above, self-efficacy — derived from social cognitive theory — positively influences effort, persistence, and the strategic utilization of learning tools [19]. In the context of educational AI applications, its impact is generally incremental and reliant upon metacognitive regulation [16]. Similarly, critical or socio-emotional self-competence is associated with ethical awareness, vigilance against biases, and the responsible use of AI [8], [15]. Nevertheless, the literature indicates that once metacognitive factors are accounted for, this dimension may exhibit ambivalent effects on productive outcomes such as creative engagement with AI [14], [18].

In this context, critical or socio-emotional self-competence is recognized as the capacity to adopt a reflective, ethical, and emotionally regulated approach concerning the use of intelligent systems [5], [10]. This facet encompasses awareness of potential risks and biases, contemplation of ethical considerations and the limitations inherent in AI applications, as well as the willingness to scrutinize outcomes and assume accountability for decisions supported by these technologies [10], [14]. Consequently, this competence facilitates the positioning of the student as an active and responsible participant in their engagement with AI systems, extending beyond mere instrumental use.

Despite its significance, recent empirical evidence indicates that the relationship between critical self-competence and productive outcomes can be complex. Elevated critical vigilance may raise the threshold for action, thereby fostering a more cautious and reflective approach to the creative application of artificial intelligence. In educational settings, it has been observed that, when controlling for metacognition and self-efficacy, this dimension does not consistently contribute additional explanatory variance and may even exhibit a negative correlation with AI-assisted creation. This reflects an underlying tension between immediate productivity and ethical responsibility in the use of generative technologies.

Although critical literacy in AI and its skills offer a promising framework, its development faces tensions and obstacles. Studies warn that without regulation and pedagogical strategies, the use of AI can promote patterns that erode the metacognitive, self-regulatory, and critical capacities that are intended to be cultivated. It is essential to examine the cognitive risks of AI-mediated learning and the conditions that limit its active deployment.

Cognitive risks in AI-mediated learning

Multiple investigations have cautioned that the deployment of artificial intelligence systems within educational environments may facilitate processes of cognitive offloading, particularly when students consistently assign tasks such as reasoning, verification, or decision-making to the technology [14]. Although this phenomenon can alleviate initial cognitive load, it also entails risks linked to overreliance, diminished metacognitive regulation, and reduced depth of learning in the absence of conscious management of AI use [15]. Specifically, persistent delegation of vital processes can impair the monitoring of comprehension and critical assessment of response quality, thereby impacting the development of transferable knowledge.

Research in cognitive psychology has shown that overconfidence in one's own performance is a critical factor that negatively affects learning, as it leads to an incorrect assessment of knowledge and the adoption of ineffective strategies [16]. In AI-mediated environments, this overconfidence can be amplified by the fluency and apparent correctness of the generated responses, reinforcing an uncritical acceptance of the information and reducing the willingness to compare it with external evidence [17]. In this sense, the formative integration of AI does not solely depend on access to tools but on the student's ability to decide when, how, and why to use them, maintaining active control over their own learning process [19], [20].

From the perspective of self-regulated learning, it has been shown that students who actively monitor their understanding, evaluate the quality of sources, and adjust their strategies achieve better learning outcomes, even in complex technological environments [18]. Recent empirical studies in higher education indicate that the effective integration of AI in learning largely depends on students' ability to regulate their interaction with technology, avoiding excessive delegation of key cognitive processes [20]. Likewise, systematic reviews have noted that much of the research on AI applied to education has focused on technological development, often overlooking the role of the student as an active agent in learning and the need for pedagogical frameworks that promote critical thinking and metacognition [21]. These findings reinforce the importance of addressing AI literacy by explicitly integrating components of self-regulation and metacognitive development, especially in formative experiences aimed at solving complex problems.

Studies demonstrate the potential and limitations of artificial intelligence in the educational field, particularly in the absence of effective regulation. Factors such as metacognition, self-efficacy, and critical self-competence are fundamental, although their impact is conditioned by circumstances and environments. It is essential to develop an integrated theoretical framework that allows for understanding the role of artificial intelligence in literacy related to this technology and guiding future research.

Theoretical integration and hypothesis formulation

Overall, the reviewed literature supports an integrated approach that articulates metacognition, self-efficacy, and critical self-competence to explain AI literacy in university contexts [1], [2], [5], [6], [15], [18], [22]. From this perspective, metacognition emerges as the most proximal determinant of effective use and knowledge of AI, as it structures the planning, monitoring, and regulation of interactions with intelligent systems. Self-efficacy contributes an incremental motivational component that fosters exploration and persistence, while critical self-competence plays a modulatory role, aimed at ensuring an ethical, reflective, and responsible use of technology.

Based on this theoretical framework, the following hypotheses are proposed:

- H1. Metacognitive awareness in the use of AI is positively associated with key dimensions of AI literacy.
- H2. Self-efficacy in AI is positively associated with AI literacy, beyond metacognitive awareness.
- H3. Critical self-competence in AI is associated with AI literacy, and its contribution may differ across dimensions, particularly those involving productive use (e.g., creation with AI).

These hypotheses guide the analytical model of the study and facilitate the empirical evaluation of the hierarchical and differential role of cognitive and socio-emotional factors in the educational use of artificial intelligence.

Methodology

A cross-sectional quantitative study was conducted. Data collection was carried out using a self-administered questionnaire with Likert-scale items.

The sample included 129 undergraduate students from an engineering school at a private university in Chile, enrolled in programming and database courses. The instrument was administered in the classroom with informed consent. The average age was 24.84 years ($SD = 6.50$), ranging from 18 to 44 years. Most participants identified as male (84.4%; $n = 108$), while 15.6% identified as female ($n = 20$). In terms of age range, 37.5% were 20 years old or younger ($n = 48$), 43.0% were between 21 and 30 years old ($n = 55$), and 19.5% were over 30 years old ($n = 25$).

Data collection was carried out through a self-administered questionnaire, structured into two sections, aimed at characterizing competencies and beliefs related to the use of artificial intelligence (AI) in learning contexts. The first section corresponded to the Meta AI Literacy Scale (MAILS) [2], which assesses different facets of AI literacy, including knowledge and conceptual understanding, use and application in tasks, detection of AI-mediated or generated content, and ethical considerations. Additionally, it incorporates dimensions of self-efficacy for solving tasks with AI, socio-emotional components related to self-regulation (e.g., emotional regulation and literacy in persuasion), and notions associated with the creation or development of AI-based solutions. The second section involved the Metacognitive Awareness in AI Use (MAI-U) instrument, adapted from [1], which measures metacognitive awareness related to AI

use, encompassing processes of planning, monitoring, and regulation of performance, and integrating technics components (e.g., system limits, reliability assessment) and ethical components (e.g., biases, transparency, and responsibility). For the analysis, the MAI-U was operationalized as a total score, while MAI-S was analyzed through its subdimensions.

The data analysis was organized into three stages: descriptive, correlational, and inferential. First, descriptive statistics (mean, standard deviation, skewness, and kurtosis) were calculated to characterize the scores of the dimensions and to assess the shape of the distributions. Then, Pearson correlations between constructs were estimated to verify the expected pattern of associations and to support the specification of predictive models. Given the overlap among facets of the MAI-U, it was operationalized as a global score (MAI_global or first component) to reduce multicollinearity and promote parsimony.

Finally, hierarchical block regressions were applied, first incorporating sociodemographic controls (e.g., gender and age range), then MAI_global as a proximal regulator, followed by AI Self-efficacy and AI Self-competence. The models were justified based on the outcome: (i) Using and applying AI, with a central role for MAI_global and additional contributions from self-efficacy and self-competence; (ii) Knowing and understanding AI, mainly associated with MAI_global and to a lesser extent with self-efficacy; (iii) Detecting AI, where it is justifiable to include self-competence before or alongside self-efficacy due to its link to critical vigilance; and (iv) Creating AI, with greater dependence on the technics-cognitive component of MAI_global and potentially modest contributions from self-efficacy. To avoid tautology, other subscales of AI literacy were not included as predictors when the dependent variable was a sub-dimension of the same domain.

At each step of the hierarchical models, the incremental contribution of the blocks was evaluated through the change in explained variance (ΔR^2), along with fit and accuracy indicators such as RMSE and information criteria (AIC and BIC). The interpretation of the effects was based on standardized coefficients (β), using a significance level of $\alpha = .05$. Additionally, the model assumptions were verified through multicollinearity diagnostics (tolerance and VIF) and influence statistics to identify potential influential observations.

Results

Table 1 presents the descriptive statistics of the instrument's dimensions and subdimensions (mean, standard deviation, skewness, and kurtosis). Overall, the average scores ranged from *moderate to high levels* (M between 3.295 and 4.219), with greater agreement in *LIT know & understand* ($M = 4.219$, $SD = 0.583$), *AI Self-efficacy* ($M = 4.156$, $SD = 0.627$), and *AI Self-competence* ($M = 4.076$, $SD = 0.663$), while *Create* recorded the lowest average ($M = 3.295$, $SD = 0.795$), suggesting a lower self-perception regarding competencies related to the creation/development of AI. The dispersion was moderate overall (SD between 0.529 and 0.850), with greater variability observed in *LIT ethics* ($SD = 0.850$) and *Create* ($SD = 0.795$), and less variability in *Ethics* ($SD = 0.529$) and in the overall *MAI index*, collapsed into a single dimension (*MAI, AI metacognition*; $M = 3.756$, $SD = 0.538$). Regarding the shape of the distributions, skewness was negative across all dimensions (skew between -0.044 and -

1.124), indicating a tendency of responses toward higher values; most skewness and kurtosis indicators remained within ranges compatible with *approximate* normality (around $|\text{skew}| < 1$ and $|\text{kurtosis}| < 1$), although *LIT know & understand* showed a more pronounced negative skewness ($\text{skew} = -1.124$), consistent with a higher concentration of high scores. Overall, these results describe a favorable response profile, with limited heterogeneity and distributions mostly close to normality, considering that *MAI* was operationalized as a unidimensional composite score for practical purposes.

Table 1. Descriptive statistics for study dimensions

Dimensions	Mean	Std. Deviation	Skewness	Kurtosis
AI Self-competence	4.076	0.663	-0.739	0.553
AI Self-efficacy	4.156	0.627	-0.808	0.509
Cognition	3.943	0.613	-0.422	0.270
Create	3.295	0.795	-0.044	-0.331
Ethics	3.501	0.529	-0.473	-0.355
Technics	3.825	0.612	-0.181	-0.028
LIT Use & apply	3.990	0.595	-0.131	-0.811
LIT know & understand	4.219	0.583	-1.124	0.826
LIT detect	3.703	0.722	-0.218	-0.008
LIT ethics	3.799	0.850	-0.606	0.167
MAI AI metacognition	3.756	0.538	-0.265	-0.399

Table 2 presents the results of Pearson correlations. A consistent pattern of positive associations can be observed, mostly moderate to high, among the constructs, which supports the relevance of moving forward with hierarchical regression models. First, the overall *MAI* index shows particularly high associations with dimensions close to its domain (*Cognition*, *Ethics*, and *Technics*; $r = .920, .917$, and $.921$, respectively; $p < .001$), suggesting significant conceptual overlap and, from a methodological standpoint, recommending operationalizing *MAI* as a single composite score (or as the first component) to reduce multicollinearity and avoid redundancies when modeling.

In line with the theoretical framework, *Using and Applying AI* (*LIT Use & apply*) is strongly associated with *MAI* ($r = .763, p < .001$) and also with *AI Self-efficacy* ($r = .696, p < .001$) and *AI Self-competence* ($r = .655, p < .001$), which supports the assumption that applied performance is linked both to metacognitive regulation/planning and to self-efficacy beliefs and self-regulatory competencies. For *Knowing and Understanding AI* (*LIT know & understand*), the association with *MAI* is equally substantial ($r = .681, p < .001$), and it is complemented by moderate correlations with *AI Self-efficacy* ($r = .605, p < .001$) and *AI Self-competence* ($r = .593, p < .001$); this pattern is consistent with the idea that the conceptual component benefits from metacognitive study strategies, while efficacy beliefs contribute secondarily.

In the case of *Detect AI* (*LIT detect*), the associations are more moderate but show a relevant nuance regarding the order of predictor entry: the correlation with *MAI* is significant ($r = .568, p < .001$), and the association with *AI Self-competence* ($r = .505, p < .001$) exceeds that observed with *AI Self-efficacy* ($r = .354, p < .001$). Since *AI Self-competence* incorporates components such as persuasion literacy and emotion

regulation, this pattern is consistent with the idea that detecting generation/mediation and evaluating outcomes requires not only metacognitive monitoring but also critical vigilance and skills to sustain judgment in the face of influence attempts, making it justifiable to include *AI Self-competence* before or alongside *AI Self-efficacy* in the final stage of the model for this dependent variable.

Table 2. Pearson correlations between the instrument dimensions and the MAI global index.

Variable	1	2	3	4	5	6	7	8	9	10	11
1. AI Self-competence	— —										
2. AI Self-efficacy	0.544 < .001	— —									
3. Cognition	0.723 < .001	0.627 < .001	— —								
4. Create	0.280 .001	0.375 < .001	0.449 < .001	— —							
5. Ethics	0.692 < .001	0.580 < .001	0.774 < .001	0.447 < .001	— —						
6. Technics	0.661 < .001	0.613 < .001	0.753 < .001	0.591 < .001	0.778 < .001	— —					
7. LIT Use & apply	0.655 < .001	0.696 < .001	0.689 < .001	0.536 < .001	0.652 < .001	0.758 < .001	— —				
8. LIT know & understand	0.593 < .001	0.605 < .001	0.608 < .001	0.313 < .001	0.656 < .001	0.618 < .001	0.631 < .001	— —			
9. LIT detect	0.505 < .001	0.354 < .001	0.445 < .001	0.470 < .001	0.500 < .001	0.619 < .001	0.459 < .001	0.378 < .001	— —		
10. LIT ethics	0.424 < .001	0.207 .019	0.513 < .001	0.434 < .001	0.634 < .001	0.492 < .001	0.347 < .001	0.422 < .001	0.362 < .001	— —	
11. MAI	0.753 < .001	0.661 < .001	0.920 < .001	0.541 < .001	0.917 < .001	0.921 < .001	0.763 < .001	0.681 < .001	0.568 < .001	0.589 < .001	— —

Finally, *Create AI (Create)* shows moderate associations with *MAI* ($r = .541, p < .001$) and with related technics/cognitive dimensions (*Technics*, $r = .591$; *Cognition*, $r = .449$; $p < .001$), along with more modest relationships with *AI Self-efficacy* ($r = .375, p < .001$) and particularly with *AI Self-competence* ($r = .280, p = .001$). This profile supports the expectation that the technics–cognitive/metacognitive component is central to 'creating/automating,' while self-efficacy could add limited incremental variance; at the same time, the high correlation between *MAI* and the related constructs reinforces the need for caution due to potential multicollinearity and overlapping explanations when additional predictors are included.

According to the correlation results, four hierarchical linear regressions were estimated with robust HC3 errors and collinearity diagnostics. In all cases, Step 0 included the controls (gender and age ranges), Step 1 added the global metacognition index in AI (*MAI*), Step 2 incorporated self-efficacy in AI, and Step 3 included self-competence in AI.

For Use & Apply AI (Table 3 and Table 4), the control block was not significant ($R^2 = .037$, $p = .192$). When incorporating *MAI*, a substantial increase in explained variance was observed ($\Delta R^2 = .554$; $R^2 = .591$, $p < .001$), with a strong and positive effect of *MAI* ($B = 0.842$, $SE = 0.065$, 95% CI [0.713, 0.971], $\beta = .761$, $p < .001$). The addition of self-efficacy contributed incremental variance ($\Delta R^2 = .070$; $R^2 = .662$, $p < .001$), remaining significant ($B = 0.337$, $SE = 0.067$, 95% CI [0.205, 0.470], $\beta = .335$, $p < .001$) while *MAI* continued as the dominant predictor ($B = 0.584$, $SE = 0.078$, 95% CI [0.429, 0.740], $\beta = .528$, $p < .001$). Self-competence did not improve the fit ($\Delta R^2 = .008$, $p = .082$) and did not show a unique effect ($B = 0.120$, $SE = 0.069$, $\beta = .140$, $p = .082$).

Table 3. Hierarchical regression model fit for predicting Use & Apply AI

Model	R ²	Adjusted R ²	R ² Change	F Change	df1	df2	p
M ₀	0.037	0.014	0.037	1.604	3	124	.192
M ₁	0.591	0.578	0.554	166.650	1	123	< .001
M ₂	0.662	0.648	0.070	25.410	1	122	< .001
M ₃	0.670	0.654	0.008	3.070	1	121	.082

Note. *M*₀ includes Gender, Age range; *M*₁ includes Gender, Age range, *MAI*; *M*₂ includes Gender, Age range, *MAI*, AI Self-efficacy; *M*₃ includes Gender, Age range, *MAI*, AI Self-efficacy, AI Self-competence

Table 4. Final hierarchical regression model coefficients for Use & Apply AI

DV		beta	t	p	95% CI		Collinearity Statistics	
					Lower	Upper	Tolerance	VIF
Use & Apply	(Intercept)		0.831	.408	-0.287	0.703		
	Gender (Male)		1.854	.066	-0.011	0.329	0.993	1.008
	Age range (21–30)		0.234	0.816	-0.124	0.157	0.985	1.015
	Age range (31+)		-0.711	0.479	-0.236	0.111		
	<i>MAI</i>	0.430	4.777	< 0.001	0.278	0.673	0.580	1.723
	AI Self-efficacy	0.343	4.881	< 0.001	0.194	0.458	0.742	1.347
	AI Self-competence	0.140	1.752	0.082	-0.016	0.268	0.652	1.533

It can be observed in Tables 5 and 6 that for *Know & Understand AI*, *M*₀ (controls: gender, age ranges) was significant but small, $R^2 = .103$, $p = .004$. The older age groups show more conceptual knowledge (positive coefficients). When adding *MAI*, there was a substantial increase in explained variance ($\Delta R^2 = .398$; $R^2 = .502$, $p < .001$), with a strong effect of *MAI* ($B = 0.700$, $SE = 0.071$, 95% CI [0.589, 0.812], $\beta = .645$, $p < .001$). Self-efficacy added additional variance ($\Delta R^2 = .038$; $R^2 = .539$, $p = .002$) and remained significant in the combined model ($B = 0.242$, $SE = 0.076$, 95% CI [0.091, 0.394], $\beta = .261$, $p = .003$), while *MAI* continued to be the main predictor ($B = 0.515$, $SE = 0.090$,

95% CI [0.337, 0.693], $\beta = .475$, $p < .001$). Self-competence contributed a marginal/non-significant increase ($\Delta R^2 = .014$, $p = .052$; $B = 0.161$, $SE = 0.082$, 95% CI [-0.001, 0.323], $\beta = .183$).

Table 5. Hierarchical regression model fit for predicting Know & Understand AI

Model	R ²	Adjusted R ²	R ² Change	F Change	df1	df2	p
M ₀	0.103	0.082	0.103	4.757	3	124	.004
M ₁	0.502	0.485	0.398	98.302	1	123	< .001
M ₂	0.539	0.520	0.038	10.014	1	122	.002
M ₃	0.554	0.531	0.014	3.863	1	121	.052

Table 6. Final hierarchical regression model coefficients for Know & Understand AI

					95% CI		Collinearity Statistics	
DV		beta	t	p	Lower	Upper	Tolerance	VIF
Know & Understand	(Intercept)		3.797	< .001	0.519	1.648		
	Gender (Male)		-0.202	.840	-0.214	0.174	0.993	1.008
	Age range (21–30)		2.367	.020	0.031	0.351	0.985	1.015
	Age range (31+)		2.893	.005	0.091	0.487		
	MAI	0.347	3.311	.001	0.151	0.601	0.580	1.723
	AI Self-efficacy	0.244	2.987	.003	0.077	0.378	0.742	1.347
	AI Self-competence	0.183	1.966	.052	-0.001	0.323	0.652	1.533

For *Detect AI*, the control block was irrelevant ($R^2 = .005$, $p = .899$). The inclusion of *MAI* significantly increased the explained variance ($\Delta R^2 = .352$; $R^2 = .357$, $p < .001$), emerging as the only strong predictor ($B = 0.816$, $SE = 0.099$, 95% CI [0.619, 1.013], $\beta = .607$, $p < .001$). Self-efficacy did not add variance ($\Delta R^2 \approx .000$, $p = .794$; $B = -0.029$, $SE = 0.112$, $\beta = -.025$, $p = .794$), and self-competence produced a small and non-significant increase ($\Delta R^2 = .012$, $p = .131$; $B = 0.140$, $SE = 0.092$, $\beta = .168$, $p = .131$). Additionally, a modest decrease was observed in the 21–30 age group compared to the reference category (Tables 7 and 8).

For *Create AI* (Tables 9 and 10), the controls did not explain the variance ($R^2 = .033$, $p = .242$). The addition of *MAI* resulted in a substantial increase ($\Delta R^2 = .290$; $R^2 = .323$, $p < .001$), with a strong and positive coefficient for *MAI* ($B = 0.813$, $SE = 0.112$, 95% CI [0.591, 1.035], $\beta = .550$, $p < .001$). Self-efficacy did not improve the model ($\Delta R^2 = .001$, $p = .759$; $B = 0.039$, $SE = 0.126$, $\beta = .031$, $p = .759$). Instead, self-competence modestly and significantly increased the fit ($\Delta R^2 = .035$, $p = .011$), with a negative coefficient ($B = -0.345$, $SE = 0.134$, 95% CI [-0.610, -0.081], $\beta = -.288$, $p = .011$), while *MAI* remained positive and robust ($B = 1.082$, $SE = 0.185$, 95% CI [0.715, 1.449], $\beta = .529$, $p < .001$). This pattern suggests that, given the same level of metacognitive regulation, higher levels of vigilance and/or socio-emotional self-regulation are associated with a more cautious approach to AI-assisted creation.

Table 7. Hierarchical regression model fit for predicting Detect AI

Model	R ²	Adjusted R ²	R ² Change	F Change	df1	df2	p
M ₀	0.005	-0.019	0.005	0.196	3	124	.899
M ₁	0.357	0.336	0.352	67.449	1	123	< .001
M ₂	0.358	0.331	0.000	0.068	1	122	.794
M ₃	0.370	0.338	0.012	2.316	1	121	.131

Table 8. Final hierarchical regression model coefficients for Detect AI

DV		beta	t	p	95% CI		Collinearity Statistics	
					Lower	Upper	Tolerance	VIF
Detect	(Intercept)		1.837	.069	-0.060	1.602		
	Gender (Male)		-0.066	.947	-0.295	0.276	0.993	1.008
	Age range (21–30)		-2.348	.021	-0.515	-0.044	0.985	1.015
	Age range (31+)		-1.674	.097	-0.538	0.045		
	MAI	0.506	4.066	< .001	0.349	1.011	0.580	1.723
	AI Self-efficacy	-0.040	-0.415	.679	-0.268	0.175	0.742	1.347
	AI Self-competence	0.168	1.522	.131	-0.055	0.422	0.652	1.533

Collinearity was assessed at all steps using VIF and tolerance. In the final models, VIF values ranged from 1.01 to 1.72 and tolerances from .58 to .99, well below the usual concern thresholds ($VIF \geq 5$ or $tolerance \leq .20$). Consequently, the estimated unique effects for MAI, self-efficacy, and self-competence are not biased by problematic collinearity.

Table 9. Hierarchical regression model fit for predicting Create AI

Model	R ²	Adjusted R ²	R ² Change	F Change	df1	df2	p
M ₀	0.033	0.010	0.033	1.414	3	124	.242
M ₁	0.323	0.301	0.290	52.581	1	123	< .001
M ₂	0.323	0.295	0.001	0.094	1	122	.759
M ₃	0.359	0.327	0.035	6.669	1	121	.011

Discussion

This study provides empirical evidence positioning metacognitive awareness in the use of AI as a central component for understanding how engineering and computing students relate to generative AI tools in core subjects [6], [8]. Descriptively, the overall profile is favorable (means between 3.295 and 4.219), with high scores in Know & Understand ($M = 4.219$), self-efficacy ($M = 4.156$), and self-competence ($M = 4.076$), while Create records the lowest average ($M = 3.295$), suggesting a lower self-perception in skills related to creation and development with AI. This asymmetry between

“understand/use” and “create” is relevant for interpreting the differential role of the predictors in the study.

Table 10. Final hierarchical regression model coefficients for Create AI

DV		beta	t	p	95% CI		Collinearity Statistics	
					Lower	Upper	Tolerance	VIF
Create	(Intercept)		0.592	.555	-0.647	1.198		
	Gender (Male)		0.621	.536	-0.217	0.416	0.993	1.008
	Age range (21–30)		-1.028	.306	-0.397	0.126	0.985	1.015
	Age range (31+)		1.255	.212	-0.118	0.528		
	MAI	0.732	5.835	< .001	0.715	1.449	0.580	1.723
	AI Self-efficacy	0.056	0.573	.568	-0.175	0.317	0.742	1.347
	AI Self-competence	-0.288	-2.582	.011	-0.610	-0.081	0.652	1.533

The consistency with which the global MAI index predicts use/application, knowledge, detection, and creation suggests that, beyond access to technology, the decisive factor is the student's ability to self-regulate their interaction with the tool: anticipating objectives, monitoring the quality of what is produced, and regulating decisions in the face of uncertainty and possible errors [1], [6], [17]. In hierarchical models, MAI explains large increases in variance when incorporated after controls: in Use & Apply, the increase is $\Delta R^2 = .554$ ($R^2 = .591$), with a strong effect ($\beta = .761$, $p < .001$); in Know & Understand, the increase is $\Delta R^2 = .398$ ($R^2 = .502$), also with a robust effect ($\beta = .645$, $p < .001$); and in Create, the increase is $\Delta R^2 = .290$ ($R^2 = .323$), with a high positive coefficient ($\beta = .550$, $p < .001$). These effect sizes empirically support the idea that metacognition functions as a proximal mechanism of self-regulated learning: when students plan the use of AI, monitor the process, and regulate strategies, they are more likely to achieve effective and transferable results [1], [2], [11], [16]. In courses such as programming and databases, where empirical validation through execution, testing, and queries allows for result verification, metacognition acts as a bridge between the support provided by AI and effective learning [6], [8], [17].

Self-efficacy showed incremental contributions mainly in instrumental dimensions such as use/application and knowledge, which aligns with sociocognitive frameworks that indicate that perceived capability facilitates persistence, exploration, and strategy adoption [2], [7], [11]. In Use & Apply, self-efficacy adds additional variance after MAI ($\Delta R^2 = .070$; $\beta = .335$, $p < .001$), remaining significant as long as MAI continues to be the dominant predictor. In Know & Understand, it also contributes incremental variance ($\Delta R^2 = .038$) and remains significant in the combined model ($\beta = .261$, $p = .003$), although once again with MAI as the main predictor ($\beta = .475$, $p < .001$). In contrast, in Detect, self-efficacy does not show a unique effect ($\beta = -.040$, $p = .679$), and in Create, it does not improve the model ($\Delta R^2 = .001$, $p = .759$). This pattern suggests that confidence without explicit regulatory processes may be insufficient, especially when AI provides plausible but incorrect responses [9], [10].

A particularly relevant result is the differential effect of critical self-competence on creation with AI. In Use & Apply, self-competence does not improve the fit ($\Delta R^2 =$

.008, $p = .082$), and in Know & Understand, its contribution is marginal/not significant ($\Delta R^2 = .014$, $p = .052$), while in Detect it does not add variance either ($p = .131$). However, in Create, it does significantly increase the fit ($\Delta R^2 = .035$, $p = .011$) but with a negative coefficient ($\beta = -.288$, $p = .011$), while MAI remains positive and robust. This pattern may reflect a more cautious use associated with greater ethical and critical awareness—sensitivity to risks, biases, attribution, and academic integrity—raising the threshold for creative action [5], [10], [14]. In this sense, critical self-competence could function as a moderating factor that tensions immediate productivity and responsibility, differentiating between “productive use” and “responsible use” of AI, especially in generation tasks [8], [15].

Regarding the controls, the effects were generally limited and did not alter the main pattern dominated by MAI, which is consistent with previous evidence reporting modest and mixed demographic associations and recommending their inclusion mainly as controls when evaluating self-regulatory predictors [24]. In Use & Apply, the control block was not significant ($R^2 = .037$, $p = .192$), and the gender coefficient was not significant in the final model. In Know & Understand, the controls were significant but explained a small amount of variance ($R^2 = .103$), and it was observed that older groups show more conceptual knowledge (positive coefficients). In Detect, the 21–30 age group shows a significant negative coefficient ($p = .021$), suggesting age-specific differences in this dimension. Overall, this is consistent with the interpretation that metacognitive regulation explains a substantial portion of performance in AI literacy beyond demographic differences.

From an applied perspective, the results reinforce that AI literacy should be approached from an integrated approach, articulating technical, critical, ethical, and metacognitive dimensions [3]–[5], [15]. For computing education, this involves designing learning experiences in which the use of AI is accompanied by explicit verification criteria (e.g., tests, decision traceability, comparison with external evidence) and guided reflection spaces on errors, biases, and attribution, promoting intellectual autonomy, critical vigilance, and transfer of learning [6], [8], [17]. Finally, the fact that the final models show low collinearity indicators (VIF between 1.01 and 1.72; tolerances between .58 and .99) supports the robustness of the interpretation of unique effects for AI self-efficacy and self-competence in the reported results.

This study presents some limitations that should be considered when interpreting the results. First, the cross-sectional, self-report design limits causal inference and may introduce social desirability biases or overestimate competencies; future studies could incorporate longitudinal designs and performance measures, such as authentic programming or database tasks with verification criteria, to compare declared literacy with observable outcomes. Second, the sample comes from a specific institutional and disciplinary context, which restricts generalization; it would be relevant to replicate the model in other universities, levels of progress, and STEM areas, also evaluating the invariance of the instrument across groups. Third, and importantly, the instrument could not be fully validated in this sample, so a larger sample size is needed to confirm its factorial structure and strengthen validity evidence (e.g., confirmatory analysis and invariance tests) before generalizing the findings to other cohorts. Fourth, the high associations between MAI and some dimensions suggest possible conceptual overlap; future work could deepen discriminant validity through CFA/SEM, explore hierarchical models, and contrast the role of metacognitive subcomponents instead of a global index.

Finally, given the differential and negative effect of critical self-competence on AI-assisted creation, further research is needed to examine underlying mechanisms such as ethical caution, fear of making mistakes, academic integrity norms, and pedagogical conditions that enable the integration of productivity and responsibility; in this vein, experimental or intervention studies could evaluate whether verification activities, guided reflection, and formative feedback simultaneously strengthen effective use, critical thinking, and responsible AI creation.

Beyond these constraints, three additional considerations regarding generalizability merit explicit discussion. First, because both predictors and outcomes were measured through self-report instruments administered in a single session, common-method variance may inflate the magnitude of the observed associations; future studies should incorporate independent or behavioral indicators (e.g., performance on AI-mediated tasks, instructor-rated artifacts) to triangulate the self-report evidence. Second, the sample is predominantly male (84.4% men, 15.6% women), which constrains the generalizability of the findings: documented gender differences in AI self-efficacy and engagement patterns suggest that the reported associations may not hold uniformly across more gender-balanced cohorts, and replication with stratified or balanced samples is recommended. Third, the institutional context—a private engineering school in Chile—likely embeds specific policies, infrastructures, and cultural norms regarding AI use in academic activities; public universities and other disciplinary contexts (e.g., humanities, design, social sciences) may show different patterns of AI literacy and metacognitive engagement, particularly where institutional AI-use policies differ in degree of permissiveness, monitoring, or pedagogical scaffolding.

Conclusions

This study aimed to analyze the relationship between metacognitive awareness in the use of artificial intelligence, self-efficacy, and critical self-competence, as well as different dimensions of AI literacy among engineering students enrolled in core courses of their program. The results allow for the extraction of relevant conclusions both at an empirical and educational level, providing evidence to better understand the educational use of AI in complex academic contexts.

First, the findings consistently show that metacognition applied to the use of AI is the most robust predictor of AI literacy outcomes. The ability to plan, monitor, and regulate interaction with AI systems is significantly associated with key dimensions such as usage and application, knowledge and understanding, error detection, and creation with AI. This pattern remains even after controlling for demographic and psychological variables, which reinforces the central role of metacognition as a proximal mechanism that shapes the quality of using these technologies in university learning. Substantively, these results suggest that the decisive factor is not just access to tools, but the ability to maintain active control over the process: defining criteria, verifying results, and adjusting strategies when uncertainty arises.

Secondly, self-efficacy in AI demonstrates an incremental contribution, particularly in instrumental dimensions. Students with higher perceived confidence tend to explore tools more actively and persist in the face of difficulties, which favors usage and application and, to a lesser extent, the associated knowledge. However, its effect remains subordinate to that of metacognition, suggesting that confidence, in the absence of explicit cognitive regulation processes, does not guarantee effective, critical, or

transferable use of AI in demanding academic tasks. Consequently, promoting self-efficacy without accompanying it with monitoring and verification strategies could be insufficient to sustain deep learning.

A third relevant finding relates to the differential role of critical self-competence. While this dimension is essential for responsible and ethical use of AI, the results indicate that it does not consistently increase productive outcomes and may be associated with a more cautious approach to creating with AI. This finding suggests a formative tension: greater ethical and critical awareness could raise the threshold for creative action, promoting more thoughtful and prudent decisions regarding the use of generative systems. Rather than interpreting this as a deficit, this pattern can be understood as a sign of critical vigilance that, under certain conditions, moderates immediate productivity to safeguard academic integrity, attribution, and risk assessment.

Overall, the results support an integrated approach to AI literacy in higher education, in which metacognition acts as a central element linking technical, motivational, and ethical skills. For engineering and computing education, these findings imply that the curricular integration of AI should not focus solely on access or efficiency but on the explicit development of metacognitive strategies that enable students to regulate their interaction with intelligent systems and sustain deep, autonomous, and responsible learning. This requires designing tasks that demand planning the use of AI, verification criteria, comparison with external evidence, and reflection on decisions, so that the tool functions as support rather than a substitute for key cognitive processes.

Finally, this study offers empirical evidence that aids in guiding curriculum development and research in AI literacy, emphasizing the significance of positioning the student as an active and reflective agent within AI-mediated educational settings.

Acknowledgments

The authors acknowledge the leadership and financial support of the School of Engineering, as well as the Educational Research and Academic Development Unit (UNIDA) for its mentorship and guidance in strengthening research capabilities in higher education.

References

- [1] B. Schraw and R. S. Dennison, "Assessing metacognitive awareness," *Contemporary Educational Psychology*, vol. 19, no. 4, pp. 460–475, 1994. doi: 10.1016/0361-476X(94)90041-8
- [2] S. Carolus, M. Weich, L. Rosenthal-von der Pütten, and N. Krämer, "Artificial intelligence literacy in higher education," *Computers and Education: Artificial Intelligence*, vol. 4, Art. no. 100126, 2023. doi: 10.1016/j.caeai.2023.100126
- [3] J. Schutte and C. Van Zyl, "Artificial intelligence literacy for students," *Computers & Education*, vol. 194, Art. no. 104698, 2023. doi: 10.1016/j.compedu.2023.104698
- [4] J. Ebert and J. Kramarczuk, "Metacognitive scaffolding for responsible AI use," *Educational Technology Research and Development*, vol. 73, no. 1, pp. 1–23, 2025. doi: 10.1007/s11423-024-10368-5

- [5] E. Kasneci, K. Sessler, S. Küchemann, M. Bannert, C. Dementieva, F. Fischer, U. Gasser, and M. Sailer, “ChatGPT for good? On opportunities and challenges of large language models for education,” *Learning and Individual Differences*, vol. 103, Art. no. 102274, 2023. doi: 10.1016/j.lindif.2023.102274
- [6] E. M. Bender, T. Gebru, A. McMillan-Major, and S. Shmitchell, “On the dangers of stochastic parrots,” *Proceedings of the ACM Conference on Fairness, Accountability, and Transparency*, pp. 610–623, 2021. doi: 10.1145/3442188.3445922
- [7] J. Dunlosky and K. J. Rawson, “Overconfidence produces underachievement,” *Learning and Instruction*, vol. 22, no. 4, pp. 271–280, 2012. doi: 10.1016/j.learninstruc.2011.08.001
- [8] T. Kosmyna, I. Lécuyer, D. Lécuyer, and F. Yang, “Cognitive offloading and overreliance on artificial intelligence,” *Nature Human Behaviour*, vol. 7, pp. 1–11, 2023. doi: 10.1038/s41562-023-01613-6
- [9] D. K. Schneider, D. Ifenthaler, and M. Yau, “Students’ regulation of learning in AI-supported environments,” *Computers & Education*, vol. 188, Art. no. 104566, 2022. doi: 10.1016/j.compedu.2022.104566
- [10] A. Zawacki-Richter, V. I. Marin, M. Bond, and F. Gouverneur, “Systematic review of research on artificial intelligence applications in higher education,” *International Journal of Educational Technology in Higher Education*, vol. 16, Art. no. 39, 2019. doi: 10.1186/s41239-019-0171-0
- [11] S. Viberg, A. Hatakka, O. Bälter, and M. Mavroudi, “The role of metacognition in AI-supported learning,” *British Journal of Educational Technology*, vol. 54, no. 2, pp. 531–547, 2023. doi: 10.1111/bjet.13298
- [12] J. Dunlosky, K. J. Rawson, E. J. Marsh, M. J. Nathan, and D. T. Willingham, “Improving students’ learning with effective learning techniques,” *Psychological Science in the Public Interest*, vol. 14, no. 1, pp. 4–58, 2013. doi: 10.1177/1529100612453266
- [13] A. Durak, F. Yilmaz, and R. Yilmaz, “Computational thinking, programming self-efficacy and problem-solving,” *Education and Information Technologies*, vol. 24, no. 5, pp. 3005–3024, 2019. doi: 10.1007/s10639-019-09905-y
- [14] S. Estaphan, A. M. McGuffin, and C. Latulipe, “Navigating AI-assisted student assignments,” *Computers and Education: Artificial Intelligence*, vol. 6, Art. no. 100211, 2025. doi: 10.1016/j.caeai.2025.100211
- [15] T. Kosmyna, I. Lécuyer, D. Lécuyer, and F. Yang, “Cognitive offloading and overreliance on artificial intelligence,” *Nature Human Behaviour*, vol. 7, pp. 1–11, 2023. doi: 10.1038/s41562-023-01613-6
- [16] S. Estaphan, A. M. McGuffin, and C. Latulipe, “Navigating the frontier of AI-assisted student assignments: Challenges, skills, and solutions,” *Computers and Education: Artificial Intelligence*, vol. 6, Art. no. 100211, 2025. doi: 10.1016/j.caeai.2025.100211
- [17] J. Dunlosky and K. J. Rawson, “Overconfidence produces underachievement: Inaccurate self-evaluations undermine students’ learning,” *Learning and Instruction*, vol. 22, no. 4, pp. 271–280, 2012. doi: 10.1016/j.learninstruc.2011.08.001
- [18] E. M. Bender, T. Gebru, A. McMillan-Major, and S. Shmitchell, “On the dangers of stochastic parrots: Can language models be too big?” *Proceedings of the ACM Conference on Fairness, Accountability, and Transparency*, pp. 610–623, 2021. doi: 10.1145/3442188.3445922

- [19] D. K. Schneider, D. Ifenthaler, and M. Yau, "Students' regulation of learning in AI-supported environments," *Computers & Education*, vol. 188, Art. no. 104566, 2022. doi: 10.1016/j.compedu.2022.104566
- [20] J. Dunlosky, K. J. Rawson, E. J. Marsh, M. J. Nathan, and D. T. Willingham, "Improving students' learning with effective learning techniques: Promising directions from cognitive and educational psychology," *Psychological Science in the Public Interest*, vol. 14, no. 1, pp. 4–58, 2013. doi: 10.1177/1529100612453266
- [21] A. Durak, F. Yılmaz, and R. Yılmaz, "Computational thinking, programming self-efficacy and problem-solving: A structural equation model," *Education and Information Technologies*, vol. 24, no. 5, pp. 3005–3024, 2019. doi: 10.1007/s10639-019-09905-y
- [22] A. Zawacki-Richter, V. I. Marín, M. Bond, and F. Gouverneur, "Systematic review of research on artificial intelligence applications in higher education," *International Journal of Educational Technology in Higher Education*, vol. 16, Art. no. 39, 2019. doi: 10.1186/s41239-019-0171-0
- [23] S. Viberg, A. Hatakka, O. Bälter, and M. Mavroudi, "The role of metacognition in learning analytics and AI-supported learning," *British Journal of Educational Technology*, vol. 54, no. 2, pp. 531–547, 2023. doi: 10.1111/bjet.13298
- [24] [24] K. Talsma, A. Chapman, and A. Matthews, "Self-regulatory and demographic predictors of grades in online and face-to-face university cohorts: A multi-group path analysis," *Br. J. Educ. Technol.*, vol. 54, no. 6, pp. 1917–1938, 2023, doi: 10.1111/bjet.13329.