

Development and Preliminary Validation of an AI-Assisted Learning Competence Scale for Ethnic Minority University Students in Yunnan, China

Yunlong Zhao¹, Shujiao Zhang², Han Zhang³, Bin Fu⁴, Yan Zhang⁵, Xian Liu⁶, Dong Yang^{7*}

^{1,2,3,4,5,6} *School of Education, Chuxiong Normal College, Chuxiong, China*

⁷ *International College, Krirk University, Bangkok, Thailand*

ABSTRACT: This study developed a self-report scale to assess AI-assisted learning competence among ethnic minority university students in Yunnan, China. Self-regulated learning and sociocultural theories guided item development within a predefined six-dimensional framework. Interviews with 16 students informed a pilot survey of 120 students, followed by confirmatory factor analysis in a main sample of 344 students. The final 18-item scale assesses learning motivation, autonomous learning, information evaluation, critical thinking, AI tool use strategies, and reflection and transfer. Subscale alpha coefficients in the pilot sample ranged from 0.771 to 0.849. The six-factor model fit well (CFI = 0.971; RMSEA = 0.053), with average variance extracted ranging from 0.552 to 0.624. Although it outperformed the single-factor model, both models showed good fit. Multigroup analyses offered preliminary support for measurement invariance, despite weaker configural fit. A review of the original interviews supported the relevance of the retained items. The scale provides a brief measure for group-level research on learning with AI, although further evidence is needed to establish the distinctness of its dimensions.

Keywords - AI-assisted learning competence, ethnic minority university students, scale development, self-regulated learning theory, sociocultural theory

1. INTRODUCTION

Learning with generative AI requires students to evaluate its responses while remaining actively engaged in their coursework. Existing frameworks identify several abilities relevant to this process. Ng et al. [1] described four components of AI literacy: knowing and understanding AI, using and applying AI, evaluating and creating AI, and AI ethics. Ridsdale et al. [2] defined data literacy as the ability to collect, manage, evaluate, and apply data critically. In China, Zhao et al. [3] used a cobweb model to visualize survey data from 2,041 students at a key provincial university, revealing substantial imbalances across dimensions of digital learning competence. Although these frameworks identify abilities relevant to learning with AI, measuring those abilities in a specific student population requires items that reflect students' coursework experiences. The present study focuses on ethnic minority students enrolled at universities in Yunnan Province. AI-assisted learning competence is defined here as students' ability to engage in learning, regulate their learning behavior, evaluate AI-generated information and reasoning, adjust their use of AI tools, and reflect on and transfer what they have learned. The scale assesses students' reports of these practices over the preceding three months. Language, culture, and access to resources are aspects of the learning context rather than separately scored dimensions.

Self-regulated learning theory and sociocultural theory guided the interpretation of the six dimensions. Zimmerman [4] emphasizes motivation, goal setting, strategy use, and metacognitive reflection, whereas Vygotsky [5] explains how social interaction and cultural tools mediate learning. These perspectives informed the examination of how students manage their learning and use AI in their everyday learning environments.

Research on hybrid human–AI regulation describes how learners and AI tools can share regulatory functions [6]. Goal setting, planning, monitoring, and strategy adjustment are also central to accounts of autonomous learning [7]. Li et al. [8] developed a measure of online learning power that encompasses motivation, adaptability, strategy use, reflection and regulation, and reciprocity. This work informed the motivation, autonomous learning, strategy use, and reflection dimensions of the present scale, while research on data literacy informed information evaluation and critical thinking [2].

The proposed scale thus covers six dimensions: learning motivation, autonomous learning, information evaluation, critical thinking, AI tool use strategies, and reflection and transfer. It focuses on reported learning practices and does not assess AI ethics, privacy, or responsibility. Accordingly, it is not intended as a comprehensive measure of AI literacy.

This study examined how a predefined six-dimensional framework could be translated into items grounded in students' learning experiences and evaluated the resulting scale's initial reliability and validity among ethnic minority university students in Yunnan.

2. METHODS

2.1 Research design

The study used a mixed-methods design to develop and evaluate a self-report scale. A literature review and semistructured interviews informed item development. The industry partner's project brief specified the six dimensions before coding began. Interview transcripts were analyzed in NVivo using open, axial, and selective coding procedures adapted from Strauss and Corbin [9]. The analysis identified initial concepts, subcategories, negative cases, and relationships within the predefined dimensions. Coding therefore elaborated the content of these dimensions rather than generating the six-dimensional framework. Interview evidence informed 48 draft items (Table 1), which were screened in a pilot survey. Confirmatory factor analysis was then used to evaluate the retained items in a separate main sample.

2.2 Participants

Purposive sampling was used to recruit 16 ethnic minority university students in Yunnan for semistructured interviews. Participants belonged to ethnic minority groups with a long-established presence in Yunnan and had some familiarity with or practical experience in AI-assisted learning or ethnic minority education. Academic advisors and corporate talent development managers helped define the project and discuss its background and construct boundaries. Only students' responses to Q01–Q14 were formally coded; contributions from advisors and managers were excluded from coding. Section 2.3 describes the item development procedures.

The pilot survey was conducted at several universities in Yunnan from July 30 to August 5, 2026. Ethnic minority students participated voluntarily through an online questionnaire. After ineligible responses, unusually rapid completions, and responses with the same option selected for every item were excluded [10], 120 valid responses remained. These data were used for item analysis, principal component analysis within each dimension, and reliability analysis.

The main survey was conducted from August 10 to August 20, 2026, at several universities in Yunnan. The online questionnaire yielded 344 valid responses from ethnic minority students who participated voluntarily. These data were used for confirmatory factor analysis, assessment of convergent validity, and multigroup comparisons across ethnic groups. Table 2 summarizes the sample characteristics.

The scale is intended for ethnic minority students enrolled at universities in Yunnan who have used generative AI for coursework within the past three months.

Participation was voluntary, and informed consent was obtained. No names, national identification numbers, or other direct identifiers were collected. Ethnicity, gender, and year of study were recorded for statistical

analysis and scale validation. Use of the data was limited to this study and instrument development. The Science and Technology Ethics Committee of Chuxiong Normal College approved the study (approval no. 2026-044).

2.3 Interview guide and scale development

Semistructured interviews [11], informed by self-regulated learning theory [4] and sociocultural theory [5], explored participants' understanding of learning competence in the AI era. Questions addressed learning motivation, autonomous learning strategies, information evaluation, critical thinking, strategies for using AI tools, and reflection and transfer. The interview guide comprised 16 questions. Q01–Q14 focused on participants' behavior and judgments when using AI for coursework and informed grounded theory coding and item development. Q15–Q16 elicited views on what the questionnaire should measure and how items should be worded. Responses to these questions informed content validity assessment and item wording but did not contribute to the 48-item pool or the reliability and validity analyses. All interviews were audio-recorded in full with participants' consent.

Verbatim transcripts were analyzed using open, axial, and selective coding. Established scales, relevant literature, and the definitions of the six dimensions guided item wording. Analysis of Q01–Q14 yielded 381 coded excerpts and 41 initial concepts. The research team selected 63 verbatim excerpts to develop 48 self-report items, with eight items per dimension. Each item was linked to at least one verified excerpt and its corresponding code.

Table 1. Examples linking the six-dimensional coding framework to pilot items

Initial concept	Subcategory	Dimension	Example pilot item
a5 AI sparks interest in learning	Interest and sustained engagement	Learning motivation	Q13: When AI helps me understand something I did not understand before, I become more interested in learning about it.
a8 Thinking independently before using AI	Learner agency	Autonomous learning	Q14: Before asking AI for help, I think through the problem or try to complete the task on my own.
a16 Checking against textbooks or authoritative sources	Verification using external sources	Information evaluation	Q21: I check information provided by AI against textbooks, class notes, or study materials.
a26 Examining the reasoning process	Checking reasoning and conditions of use	Critical thinking	Q16: I check whether the reasoning provided by AI is logical.
a35 Adjusting prompts	Prompt refinement	AI tool use strategies	Q47: When an AI response is not accurate enough, I make my original question more specific.
a37 Reviewing experiences of using AI	Reviewing the process	Reflection and transfer	Q18: After a learning session with AI, I review how I used it.

Note. Concepts and subcategories were derived from the coding results. Q numbers refer to items in the 48-item pilot scale, not to interview questions. Some example items were removed during screening.

Table 1 illustrates how coding informed item wording. For example, “thinking independently before using AI” was coded as “learner agency” and translated into an item about attempting a task before seeking AI assistance. “Examining the reasoning process” became an item about checking the logic of AI’s reasoning.

Three experts with doctorates in psychology or education reviewed the item content. Following revisions, ten ethnic minority university students in Yunnan assessed the clarity and appropriateness of the wording. Their feedback guided further refinements.

The pilot scale comprised 48 items across six dimensions. Respondents used a five-point Likert scale (1 = not at all true of me; 5 = completely true of me), with higher scores indicating greater AI-assisted learning competence [12]. Students were asked to report their use of generative AI for coursework over the past three months and told there were no right or wrong answers. Those who had not used AI during this period should be classified as nonusers rather than assigned low competence scores.

2.4 Data analysis

Pilot data ($N = 120$) were analyzed in SPSS 26.0. Item analysis examined critical ratios, item–total correlations, and Cronbach’s alpha if each item was deleted [13]. Principal component analysis (PCA) extracted one component from the eight candidate items in each dimension. Screening considered the Kaiser–Meyer–Olkin (KMO) statistic, Bartlett’s test of sphericity, component loadings, corrected item–total correlations (CITCs), and content coverage. The three retained items per dimension were reassessed using single-component PCA. This procedure screened items within predefined dimensions; CFA subsequently tested the full six-factor structure. Cronbach’s alpha was calculated for each subscale and the full shortened scale in the pilot sample [14].

Main survey data ($N = 344$) were analyzed in AMOS 24.0 using confirmatory factor analysis (CFA) with maximum likelihood estimation. Neither the pilot nor the main dataset used in the analyses contained missing values. A correlated six-factor model was fit and compared with a single-factor model [15]. Model fit was evaluated using χ^2 , CFI, TLI, RMSEA, and SRMR. Values near 0.95 for CFI and TLI, 0.06 for RMSEA, and 0.08 for SRMR served as guidelines rather than absolute cutoffs [16]. Convergent validity was evaluated using standardized loadings and average variance extracted (AVE), with 0.50 as the reference value for AVE [17]. Composite reliability (CR) was also calculated. Multigroup CFA tested configural, metric, and scalar invariance across ethnic groups by successively constraining factor loadings and item intercepts. Changes in fit were interpreted in conjunction with the fit of the baseline model [18].

The qualitative review used the original interview dataset, comprising 293 meaning units from 16 participants. These units differ from the coded excerpts counted during item development. The review focused on Q03–Q10, which addressed the six dimensions and contained 143 meaning units. Directed content analysis [19] and cross-case review compared the 18 retained items with descriptions of the corresponding behaviors. Support required a match in meaning, relevant accounts from at least two participants, and documented meaning-unit IDs and source locations. This was a content review of the original interviews, not validation in an independent sample.

3. RESULTS

3.1 Main sample characteristics

Table 2. Characteristics of the main sample ($N = 344$)

Variable	Category	<i>n</i>	%
Ethnic group	Yi	123	35.8%
	Lisu	76	22.1%
	Miao	72	20.9%
	Other ethnic groups	73	21.2%
Year of study	First year	123	35.8%
	Second year	151	43.9%
	Third year	60	17.4%
	Fourth year	10	2.9%
Gender	Men	107	31.1%
	Women	237	68.9%

Note. The “other ethnic groups” category includes all ethnic groups other than Yi, Lisu, and Miao. Most participants were first- or second-year students, and 68.9% were women.

3.2 Item analysis and selection

Item analysis used the 120 pilot responses. Within each dimension, participants were ranked by their total score on the eight candidate items. The top and bottom 27% formed high- and low-scoring groups of 32 participants each. Item means, standard deviations, critical ratios, CITCs, and Cronbach's alpha if each item was deleted were calculated. Across the 48 items, means ranged from 3.783 to 4.150, standard deviations from 0.686 to 0.946, and critical ratios from 6.400 to 13.281. All between-group differences were significant at $p < 0.001$, and CITCs ranged from 0.509 to 0.770. All items were retained for component analysis and content review.

KMO values ranged from 0.879 to 0.906 across the six dimensions, and all Bartlett's tests were significant at $p < 0.001$ (Table 3). A single component explained 49.3%–58.1% of the variance within each dimension. These results, component loadings, CITCs, and item content informed the selection of three items per dimension, yielding an 18-item scale.

Table 3. Pilot item screening and reliability of the retained items ($N = 120$)

Dimension	KMO	Bartlett's χ^2	df	p	Variance explained	Items retained	α
Learning motivation	0.906	418.253	28	< 0.001	55.7%	3	0.797
Autonomous learning	0.891	425.550	28	< 0.001	55.2%	3	0.849
Information evaluation	0.879	325.885	28	< 0.001	49.3%	3	0.791
Critical thinking	0.898	457.837	28	< 0.001	58.1%	3	0.818
AI tool use strategies	0.894	387.126	28	< 0.001	54.2%	3	0.795
Reflection and transfer	0.893	365.141	28	< 0.001	52.7%	3	0.771

Note. KMO, Bartlett's test, and variance explained refer to the eight candidate items in each dimension. Alpha values refer to the three retained items per dimension.

PCA of the retained items in the pilot sample yielded component loadings ranging from 0.809 to 0.898 (Table 4). The single-component solutions accounted for 68.6%–76.9% of the variance within the six dimensions. Cronbach's alpha ranged from 0.771 to 0.849 for the three-item subscales and was 0.955 for the full 18-item scale in the pilot sample ($N = 120$), supporting acceptable internal consistency.

Table 4. Final scale items, pilot component loadings, and main-sample factor loadings

Final item	Data ID	Dimension	Item	Pilot PCA	Main CFA
1	10	Learning motivation	When AI explains a concept in plain language, I am more willing to explore the topic further.	0.824	0.698
2	16	Learning motivation	When AI helps me understand something I did not understand before, I become more interested in learning about it.	0.867	0.775
3	17	Autonomous learning	Before asking AI for help, I think through the problem or try to complete the task on my own.	0.898	0.757
4	20	AI tool use strategies	I choose between basic and more specialized AI modes according to the complexity of the task.	0.846	0.803
5	21	Reflection and transfer	After a learning session with AI, I review how I used it.	0.818	0.714
6	23	Autonomous learning	When using AI, I ask it to suggest an approach and then complete the learning task myself.	0.876	0.790

Final item	Data ID	Dimension	Item	Pilot PCA	Main CFA
7	25	Critical thinking	I consider whether I can reach the same conclusion by following the reasoning provided by AI.	0.855	0.762
8	31	Critical thinking	I check whether the conditions required for the method used by AI are met.	0.849	0.787
9	34	Learning motivation	After learning with AI, I feel more confident about mastering the material.	0.844	0.754
10	36	Information evaluation	I check important information from AI against other sources to see whether it is accurate and reliable.	0.860	0.773
11	39	Reflection and transfer	I apply AI-assisted learning approaches that worked in one course to other courses.	0.809	0.751
12	41	Autonomous learning	When using AI to learn important material, I check my understanding as I go.	0.856	0.787

Table 4. Final scale items and loadings (continued)

Final item	Data ID	Dimension	Item	Pilot PCA	Main CFA
13	42	Information evaluation	When I doubt an AI response, I ask follow-up questions and ask it to check its answer again.	0.861	0.790
14	43	Critical thinking	I check whether AI's analysis contradicts its final conclusion.	0.867	0.750
15	44	AI tool use strategies	When faced with a complex problem, I first break it down into smaller problems.	0.847	0.807
16	45	Reflection and transfer	I turn approaches that worked well in a learning session with AI into methods I can use again.	0.858	0.782
17	48	Information evaluation	When I cannot find reliable sources to support AI-generated content, I use it only as a reference rather than use it directly.	0.809	0.739
18	50	AI tool use strategies	When an AI response is not accurate enough, I make my original question more specific.	0.834	0.759

Note. Pilot loadings were obtained from separate single-component PCAs of the three retained items in each dimension ($N = 120$). Main-sample loadings are standardized estimates from the six-factor CFA ($N = 344$). Data IDs refer to item identifiers in the data file; final item numbers refer to the 18-item scale.

3.3 Confirmatory factor analysis and measurement invariance

The main sample comprised 344 students: 123 Yi, 76 Lisu, 72 Miao, and 73 from other ethnic groups. The measurement structure was first evaluated using an 18-item correlated six-factor model, followed by configural, metric, and scalar invariance models (Table 5). The six-factor model fit well: $\chi^2 = 237.405$, $df = 120$, CFI = 0.971, TLI = 0.963, RMSEA = 0.053, and SRMR = 0.030. Standardized loadings for the 18 items ranged from 0.698 to 0.807.

Table 5. Measurement invariance across ethnic groups

Model	χ^2	df	CFI	TLI	RMSEA	SRMR	Δ CFI	Δ RMSEA
Configural	915.600	480	0.903	0.876	0.052	0.059	—	—
Metric	939.407	516	0.906	0.888	0.049	0.074	+0.003	-0.003
Scalar	979.353	552	0.905	0.894	0.048	0.074	-0.001	-0.001

The configural model yielded RMSEA = 0.052 and SRMR = 0.059, but weaker incremental fit indices (CFI = 0.903, TLI = 0.876). From the configural to the metric model, Δ CFI was +0.003 and Δ RMSEA was -0.003. From the metric to the scalar model, Δ CFI was -0.001 and Δ RMSEA was -0.001. Constraining loadings and intercepts produced little deterioration in fit. These findings provide preliminary support for measurement invariance, although the weaker configural fit limits the strength of this conclusion.

3.4 Convergent validity and competing models

Table 6. Convergent validity in the main sample

Dimension	Standardized loadings	Composite reliability (CR)	Average variance extracted (AVE)
Learning motivation	0.698–0.775	0.787	0.552
Autonomous learning	0.757–0.790	0.822	0.606
Information evaluation	0.739–0.790	0.811	0.589
Critical thinking	0.750–0.787	0.810	0.588
AI tool use strategies	0.759–0.807	0.833	0.624
Reflection and transfer	0.714–0.782	0.793	0.562

As shown in Table 6, standardized loadings ranged from 0.698 to 0.807, composite reliability (CR) from 0.787 to 0.833, and average variance extracted (AVE) from 0.552 to 0.624. All AVE values exceeded 0.50, supporting convergent validity within the six-factor model [17].

Table 7. Comparison of competing models

Model	χ^2	df	CFI	TLI	RMSEA	SRMR
18-item single-factor model	292.189	135	0.962	0.957	0.058	0.033
18-item correlated six-factor model	237.405	120	0.971	0.963	0.053	0.030

As shown in Table 7, the six-factor model fit better than the single-factor model ($\Delta\chi^2 = 54.784$, $\Delta df = 15$, $p < 0.001$). CFI increased from 0.962 to 0.971, and RMSEA decreased from 0.058 to 0.053. Both models nevertheless showed good fit, so the comparison alone does not establish that the six dimensions are empirically distinct.

3.5 Review of qualitative evidence

Table 8. Qualitative evidence supporting the six dimensions

Dimension	Final items	Core qualitative concepts	Meaning units	Participants
Learning motivation	1, 2, 9	Plain-language explanations make content easier to understand and encourage further learning; understanding previously unfamiliar material increases interest; AI-assisted understanding and error correction build confidence in mastering the material.	6	6

Dimension	Final items	Core qualitative concepts	Meaning units	Participants
Autonomous learning	3, 6, 12	Thinking or attempting a task independently before using AI to check or supplement the work; asking AI for an approach while completing the task oneself; checking understanding by restating ideas, solving problems independently, or attempting variations of a problem.	6	4
Information evaluation	10, 13, 17	Cross-checking important information against textbooks, research literature, or official sources; asking follow-up questions and requesting a recheck when answers are doubtful; treating AI content only as a reference when reliable supporting sources are unavailable.	7	7
Critical thinking	7, 8, 14	Independently judging whether AI's reasoning leads to the same conclusion; checking the conditions under which a method, formula, or line of reasoning applies; checking consistency between the analysis and the final conclusion.	6	6
AI tool use strategies	4, 15, 18	Choosing a basic or specialized mode according to task complexity; dividing a complex problem into smaller problems to solve step by step; rewriting broad questions to make them more specific and clearly bounded.	6	3
Reflection and transfer	5, 11, 16	Reviewing how AI was used after a learning session and drawing lessons from the experience; transferring effective AI-assisted learning approaches from one course to another; turning useful practices into methods that can be used again.	6	5

The 18 retained items were supported by 37 distinct interview meaning units representing all 16 participants (Table 8). Each item corresponded to descriptions of the relevant behavior from at least two participants. The retained items therefore remained grounded in the experiences documented during item development.

4. DISCUSSION

4.1 Interpretation of the six dimensions

Interviews with ethnic minority university students in Yunnan helped translate the predefined six-dimensional framework into 18 items describing everyday coursework practices. The items address attempting tasks independently, checking AI-generated information and reasoning, refining questions, and applying useful approaches to subsequent tasks. Reviews by students and experts helped ensure that the items described recognizable behaviors in clear language.

Consistent with self-regulated learning theory [4], the scale emphasizes motivation, independent work, monitoring, and reflection. Sociocultural theory helps explain AI's role as a tool in students' learning environments [5]. Interview accounts and feedback on wording shaped the scale for this population. The scale does not directly assess language, cultural practices, or resource access, and the findings do not establish that these behaviors are unique to ethnic minority students.

The item content distinguishes information evaluation from critical thinking. Information evaluation involves checking external sources, following up on questionable answers, and avoiding direct use of unsupported

content. Critical thinking involves examining reasoning, identifying when a method applies, and checking whether an analysis is consistent with its conclusion.

AI tool use strategies involve selecting an AI mode, breaking down complex problems, and refining questions. Autonomous learning emphasizes independent attempts, responsibility for task completion, and monitoring one's understanding. Reflection and transfer focus on reviewing experience and applying useful approaches in subsequent coursework. These distinctions capture different aspects of how students manage learning with AI.

4.2 Measurement properties and applications

The pilot subscales showed acceptable internal consistency, and the main-sample CFA, composite reliability, and AVE supported the proposed measurement structure. However, the single-factor model also fit well, and the high overall alpha is consistent with substantial overlap among dimensions. Factor correlations and discriminant validity evidence are needed to establish how clearly the subscales differ. Qualitative evidence supports item relevance but does not resolve this issue. Measurement invariance remains preliminary given the weaker configural fit.

The scale can be used in exploratory, group-level surveys of reported learning practices. Individual ranking or recruitment decisions would require evidence on score interpretation, external criteria, and stability over time. Ethnic-group comparisons should account for the combined "other ethnic groups" category, which does not distinguish among its constituent groups.

4.3 Practical implications

Instructors can use the items to guide classroom discussion and observation. Activities addressing information evaluation and critical thinking could ask students to compare sources, examine AI's reasoning, and explain when a method applies. For autonomous learning and AI tool use strategies, students could attempt tasks before using AI, compare prompts, and break complex problems into smaller steps. These activities reflect the item content; their instructional effectiveness was not tested.

Students can use the items to reflect on how they check AI responses, monitor their understanding, and apply useful approaches across courses. Researchers can use the scale to describe these reported practices, with further evaluation of item interpretation and measurement properties when extending its use to other institutions or populations.

4.4 Limitations

Purposive interview sampling and voluntary participation in online surveys limit the generalizability of the findings to ethnic minority university students across Yunnan. The main sample was predominantly female and concentrated in the first two years of study. Replication should include a broader range of institutions, years of study, and ethnic groups.

Self-reported practices may differ from actual performance, and criterion validity was not assessed against observed tasks or academic outcomes. Behavioral tasks, including situational judgment items, could complement the scale. Further validation should also examine whether the six subscales assess sufficiently distinct abilities.

Test-retest reliability and sensitivity to change were not examined. Repeated assessments could establish score stability and assess changes associated with coursework experience and growing familiarity with AI tools.

5. CONCLUSION

This study developed an 18-item self-report scale assessing six aspects of AI-assisted learning competence among ethnic minority university students in Yunnan. Interview findings, pilot reliability estimates, and CFA results from the main sample provide initial support for its content and measurement structure. The scale offers a brief instrument for group-level research and identifies learning practices that can guide classroom reflection. Further validation should clarify subscale distinctness, comparability across ethnic groups, and correspondence with observed performance.

Acknowledgments. This article is based on research conducted jointly by faculty members and students for the 2026 Houcan Cup competition.

Conflicts of interest. The authors declare no conflicts of interest.

6. REFERENCES

- [1] D. T. K. Ng, J. K. L. Leung, S. K. W. Chu, and M. S. Qiao, Conceptualizing AI literacy: An exploratory review, *Computers and Education: Artificial Intelligence*, 2, 2021, 100041. <https://doi.org/10.1016/j.caeai.2021.100041>.
- [2] C. Ridsdale, J. Rothwell, M. Smit, H. Ali-Hassan, M. Bliemel, D. Irvine, D. Kelley, S. Matwin, and B. Wuetherick, *Strategies and best practices for data literacy education: Knowledge synthesis report* (Halifax: Dalhousie University, 2015). <https://hdl.handle.net/10222/64578>.
- [3] L. Zhao, X. T. Zhang, and J. H. Liu, How to improve the digital learning ability of college students: An analysis of the spider web model based on survey data from a provincial key university, *Distance Education in China*, 45(5), 2025, 55–74 (in Chinese).
- [4] B. J. Zimmerman, Attaining self-regulation: A social cognitive perspective, in M. Boekaerts, P. R. Pintrich, and M. Zeidner (Eds.), *Handbook of self-regulation* (San Diego, CA: Academic Press, 2000) 13–39. <https://doi.org/10.1016/B978-012109890-2/50031-7>.
- [5] L. S. Vygotsky, *Mind in society: The development of higher psychological processes*, M. Cole, V. John-Steiner, S. Scribner, and E. Souberman (Eds.) (Cambridge, MA: Harvard University Press, 1978).
- [6] I. Molenaar, The concept of hybrid human-AI regulation: Exemplifying how to support young learners' self-regulated learning, *Computers and Education: Artificial Intelligence*, 3, 2022, 100070. <https://doi.org/10.1016/j.caeai.2022.100070>.
- [7] W. G. Pang, *Autonomous learning: Principles and strategies of learning and teaching* (Shanghai: East China Normal University Press, 2003) (in Chinese).
- [8] B. M. Li, L. L. Gong, and Z. T. Zhu, Development and verification of online learning power assessment tools for online learners, *Open Education Research*, 24(3), 2018, 77–84, 120 (in Chinese). <https://doi.org/10.13966/j.cnki.kfjyyj.2018.03.009>.
- [9] A. Strauss and J. Corbin, *Basics of qualitative research: Techniques and procedures for developing grounded theory*, 2nd ed. (Thousand Oaks, CA: Sage Publications, 1998).
- [10] A. W. Meade and S. B. Craig, Identifying careless responses in survey data, *Psychological Methods*, 17(3), 2012, 437–455. <https://doi.org/10.1037/a0028085>.
- [11] S. Brinkmann and S. Kvale, *InterViews: Learning the craft of qualitative research interviewing*, 3rd ed. (Thousand Oaks, CA: Sage Publications, 2015).
- [12] R. Likert, A technique for the measurement of attitudes, *Archives of Psychology*, 22(140), 1932, 1–55.
- [13] M. L. Wu, *Practice of questionnaire statistical analysis: SPSS operation and application* (Chongqing: Chongqing University Press, 2010) (in Chinese).
- [14] L. J. Cronbach, Coefficient alpha and the internal structure of tests, *Psychometrika*, 16(3), 1951, 297–334. <https://doi.org/10.1007/BF02310555>.
- [15] B. M. Byrne, *Structural equation modeling with AMOS: Basic concepts, applications, and programming*, 2nd ed. (New York: Routledge, 2010).
- [16] L.-T. Hu and P. M. Bentler, Cutoff criteria for fit indexes in covariance structure analysis: Conventional criteria versus new alternatives, *Structural Equation Modeling: A Multidisciplinary Journal*, 6(1), 1999, 1–55. <https://doi.org/10.1080/10705519909540118>.
- [17] C. Fornell and D. F. Larcker, Evaluating structural equation models with unobservable variables and measurement error, *Journal of Marketing Research*, 18(1), 1981, 39–50. <https://doi.org/10.1177/002224378101800104>.
- [18] F. F. Chen, Sensitivity of goodness of fit indexes to lack of measurement invariance, *Structural Equation Modeling: A Multidisciplinary Journal*, 14(3), 2007, 464–504. <https://doi.org/10.1080/10705510701301834>.
- [19] H.-F. Hsieh and S. E. Shannon, Three approaches to qualitative content analysis, *Qualitative Health Research*, 15(9), 2005, 1277–1288. <https://doi.org/10.1177/1049732305276687>.

INFO

Corresponding Author: [Dong Yang](#), International College, Krirk University, Bangkok, Thailand.

How to cite/reference this article: [Yunlong Zhao](#), [Shujiao Zhang](#), [Han Zhang](#), [Bin Fu](#), [Yan Zhang](#), [Xian Liu](#), [Dong Yang](#), Development and Preliminary Validation of an AI-Assisted Learning Competence Scale for Ethnic Minority University Students in Yunnan, China, *Asian. Jour. Social. Scie. Mgmt. Tech.* 2026; 8(5): 27-37.