

Academic Use of AI by College Students and Self-Regulated Learning: The Moderating Effect of Usage Motivation and the Construction/Substitution Pathways

Wushi Li

School of Humanities and Social Science, The Chinese University of Hong Kong, Shenzhen, Guangdong, China
Email: wushili@link.cuhk.edu.cn

Abstract—Generative Artificial Intelligence (AI) permeates students' educational writing and assignment conduct, but its scholarly effects emerge less from use quantity as from qualitatively different strategic purpose types: enhancing (scaffolding planning, monitoring, reflective amendment) versus replacing (subsuming central cognitive labor). This paper sketches and empirically validates an integrative framework that associates two motivational dispositions (Enhancement Orientation vs. Replacing Orientation) and three AI use patterns (A, B, C) with effects on two outcomes: 1) satisfaction and 2) quality, thereby contrasting use frequency with strategic purpose. Key constructs include motivational orientations, Self-Regulated Learning (SRL), self-efficacy, and the two AI use intentions. A multidisciplinary sample of 300 university students with recent experience using generative AI completed a questionnaire, which included readapted items measuring motivation, SRL, and self-efficacy, as well as novel items assessing supportive versus substitutional intentions. Confirmatory factor analyses (WLSMV) revealed high convergent and discriminant validity (standardized loadings 0.749–0.907; Average Variance Extracted (AVE) 0.598–0.642) of four constructs, indicating unique explanatory value for each of these four variables. Composite reliabilities ranged from 0.748 to 0.911. Structural model fit was good ($\chi^2(194) = 112.232, p = 1.000$; CFI = 1.000; TLI $\approx$ 1.00; RMSEA = 0.000; SRMR = 0.034). Both Enhancement and Efficiency-Avoidance Orientations increased SRL (and self-efficacy) and, both directly and indirectly, supportive intentions, and decreased substitutional intentions (to a greater extent) through SRL. Both orientations also directly increased substitutional intentions but did not meaningfully decrease SRL. Self-efficacy played a mostly protective role—it was a negative predictor of substitutional but not of supportive intention. These findings establish the two-way motivational system: a promotive enhancement pathway routed through SRL and a more direct Efficiency-Avoidance pathway that amplifies substitution dependence.

Keywords—generative Artificial Intelligence (AI), self-regulating learning, motivation for use, supportive and substitutive types, undergraduate student

I. INTRODUCTION

As discussed by Kasneci *et al.* [1] and Lund *et al.* [2], Large Language Model (LLM)-based generative Artificial Intelligence (AI) systems (e.g., ChatGPT, DeepSeek) have quickly spread into university students' academic writing, literature review, coding/debugging, data processing, language polishing, etc. workflows since the end of 2022. Such adoption patterns include low technical barriers, multidisciplinary impacts, and fragmented micro-usage (e.g., idea outlining, paragraph writing, paraphrasing, style polishing)—all shaped by current technology and policies. Traditionally quantitative measures, such as for example binary use (yes/no) or simply an aggregate frequency, are becoming increasingly unsatisfactory to account for diverse learning gains and security threats.

Building on Self-Regulated Learning (SRL) theory by Zimmerman [3] and cognitive offloading research by Risko and Gilbert [4], we distinguish two intention types rather than treating AI engagement as homogeneous. Supportive intention frames the system as a scaffold or thinking partner, as proposed by Wood *et al.* [5], aiming to enhance planning, monitoring, evaluative reflection, and iterative refinement of structure, argumentation, and linguistic clarity. Substitution intention, in contrast, orients toward outsourcing core generative or decision processes (e.g., “produce a full paragraph”, “rewrite this section” with limited critical interrogation), risking shallow processing and over reliance. These intentions can theoretically co-exist in a single task sequence, yet their dominance signals different regulatory mindsets and potential downstream consequences for capability building and academic integrity.

We emphasize two higher-order motivational orientations that influence how learners adopt an adaptive stance toward AI. As detailed by Dweck [6] and Biggs *et al.* [7], Enhancement Orientation captures growth- and quality-focused motivation: increasing knowledge,

strengthening argument coherence, and developing long-term skills. In parallel, drawing from the work of Elliot and Church [8] and Midgley *et al.* [9], Efficiency-Avoidance Orientation captures efforts to reduce time as well as cognitive effort and to avoid tedious aspects. Though not exclusive, a long-term preference toward enhancing motives should correspond to conscious involvement but efficiency avoidance should directly motivate substitutional strategies to bypass internal regulation processing.

SRL encapsulates proactive planning, strategic monitoring, and reflective adjustment cycles in academic tasks [3]. Self-efficacy, as explored by Schunk and DiBenedetto [10], influences persistence and ownership and may operate either as a promotive catalyst for constructive engagement or as a defensive barrier against maladaptive outsourcing. Conceptually, Enhancement Orientation is expected to elevate SRL and self-efficacy, channeling intentions toward supportive usage; Efficiency-Avoidance Orientation may act more directly on substitution intentions, potentially with weaker or null relationships to SRL.

This study investigates the psychological pathways of student engagement with academic AI. We first aim to reliably differentiate between students' supportive and substitution intentions for AI use. We then examine how enhancement and efficiency-avoidance orientations predict Self-Regulated Learning (SRL) and self-efficacy. Subsequently, we analyze how SRL and self-efficacy in turn predict these AI use intentions. Finally, we test an integrated model to determine whether SRL and self-efficacy mediate the effects of motivational orientations on supportive and substitution intentions.

II. LITERATURE REVIEW

Early commentaries and survey snapshots on this topic present several limitations. They have often: (a) emphasized frequency or generic purpose categories over qualitatively distinct intention constructs [1, 2]; (b) failed to integrate dual motivational dispositions with SRL and self-efficacy into a single explanatory chain; and (c) insufficiently connected AI usage intentions with academic integrity frameworks [11, 12]. Furthermore, much of the current discourse lacks actionable levers for learning that are grounded in validated psychological concepts. We argue that before advancing to the validation of behavioral logs, it is necessary to first establish a distilled, intention-level model. Such a model is a prerequisite for guiding future research and ensuring that behavioral data can be meaningfully interpreted.

This study develops and tests an intention-level structural model: Enhancement Orientation and Efficiency-Avoidance Orientation → SRL and Self-Efficacy → Supportive versus Substitution AI Usage Intentions. Because no objective interaction logs or performance outcomes were collected in the present phase, the model is explicitly constrained to psychological and intentional constructs. Behavioral enactments (e.g., iteration depth, verification prompts) are proposed as

avenues for future research rather than being framed here as measured effects.

Based on the theoretical framework and research questions outlined above, we propose the following hypotheses:

H1a: Enhancement Orientation positively predicts SRL and Self-Efficacy.

H1b: Efficiency-Avoidance Orientation negatively or non-significantly predicts SRL, and shows a weakly negative or null association with Self-Efficacy (exploratory).

H2a: SRL positively predicts Supportive Intention.

H2b: SRL negatively predicts Substitution Intention.

H3a: Self-Efficacy negatively predicts Substitution Intention.

H3b: Self-Efficacy has a null (or minimal) positive relation with Supportive Intention (anticipated “defensive” rather than promotive role).

H4: Enhancement Orientation indirectly increases Supportive and suppresses Substitution via SRL (and secondarily Self-Efficacy); Efficiency-Avoidance primarily exerts a direct positive effect on Substitution with weaker indirect pathways.

III. RESEARCH DESIGN

The target population of this study consists of undergraduate first-year to doctoral students who have already acquired basic academic writing skills and have used generative AI at least once in the most recent semester. This range takes into account both the frequency of writing tasks and the developmental plasticity of learners. We selected a total of 300 college students to fill out the questionnaire. Based on their majors, the students were classified into science and engineering, humanities and social sciences, economics and management, art and design, and other majors. The micro-contextual randomized experiments inserted in the questionnaire were randomly assigned through a computer within the online questionnaire system.

The study used a self-developed questionnaire entitled “University Students’ Academic Use and Self-Regulation of Generative AI.” For this questionnaire, items measuring constructs related to the key concepts were partially adapted from existing scales, while other content was developed *de novo* to fit the specific context and not directly reused from the original scales. Specifically: The AI usage motivation items (15–22) were contextualized in an upward/avoidance effort based on achievement goal theory and mastery/avoidance orientation [8, 9] and deep/surface learning orientation semantics [7]; The self-regulation learning items (23–28) borrowed its conceptions from metacognitive dimensions of planning, monitoring, reflection in MSLQ [13], keeping the generalization words in, but adjusting the functional context. All the artificial use frequency behaviors (7–14, differentiated between “supportive” and “substitutive” concrete behaviors) and the 3 micro-context manipulation passages (constructive passive dependence/neutral) and items (35–37), and the strategy intent items (38–44, constructive passive dependence) were also newly

developed. Thirty students were recruited to participate in cognitive interviews to ensure the semantic clarity and comprehensibility of the questionnaire items. Finally, items assessing the drivers of rewriting behavior (i.e., constructive motivation, SRL, self-efficacy, and integrity attitude constructs) were included to supplement the original contextualized behavior and strategy intention dimensions. This allows for the simultaneous capture of the dual dimensions of “AI usage type differences” and “learning regulation/ethical boundary cognition”, providing a measure of validity (both contextual and discriminant) for future mechanism modeling.

In order to examine the influence of different types of content examples (context manipulation) on the participants’ immediate strategic intentions, ability-related beliefs (especially the sense of active analysis and the sense of irrelevance), and attitudes towards the “acceptable boundaries of AI usage”, this study set up a micro-context randomized experiment. The questionnaire included a random context that randomly appeared in three scenarios. The scenarios were as follows: the supporting exemplary scenario: presenting a concise case where a student first wrote step-by-step, then compared, annotated, verified, and finally rewrote; the alternative risk scenario: presenting a case where a student directly said “please rewrite” and adopted a large amount of unfiltered content, which later was pointed out for its lack of reasoning and factual errors; the neutral control scenario: only providing the task background and “AI-assisted writing is allowed” explanation, without providing a demonstration of strategies. The readability indicators of the scenario descriptions were basically the same.

The questionnaire is stored on the Wenjuanxing platform, that satisfied the academic research data encryption and log record. An electronic informed consent appeared before participant access to questionnaire, explaining the research purpose, data anonymization, can quit the questionnaire anytime, refusing the questionnaire doesn’t have adverse consequences, researcher’s contact information is included. The participants can press only after clicking “Agree”. The platform provided pseudo-random number generates the micro-context allocation of randomization.

IV. RESULTS

An analysis was conducted using a Structural Equation Model (SEM) with the WLSMV estimator in the R programming language to test the hypothesized relationships between motivational orientations, self-regulation, self-efficacy, and AI use intentions. The results are presented following the sequence of measurement model evaluation, overall structural model fit, mediation analysis, and hypothesis testing.

A. Measurement Model Quality

The measurement model demonstrated strong quality. Convergent validity was established, as all standardized item loadings were significant and ranged from 0.749 to 0.907. Average Variance Extracted (AVE) values

(0.598–0.774) and Composite Reliability (CR) scores (0.748–0.911) surpassed their respective thresholds of 0.50 and 0.70. Internal consistency was also supported, with Cronbach’s alpha and McDonald’s omega values ranging from 0.748 to 0.910. Discriminant validity was confirmed using both the Fornell–Larcker criterion and the HTMT criterion (all values < 0.85).

B. Overall Model Fit and Structural Paths

The overall structural model demonstrated an exemplary fit to the data ($\chi^2(194) = 112.232$, $p = 1.000$; CFI = 1.000; TLI = 1.000; RMSEA = 0.000; SRMR = 0.034; WRMR = 0.556).

The standardized path coefficients are presented in the table below and reveal several key relationships.

TABLE I. STANDARDIZED PATH COEFFICIENTS

SRL	Enhance	0.646	<0.001	***
SRL	EffAvoid	−0.044	0.359	ns
SelfEff	Enhance	0.251	<0.01	**
SelfEff	EffAvoid	0.072	0.258	ns
SelfEff	SRL	0.321	<0.01	**
SupportIntent	SRL	0.291	<0.001	***
SupportIntent	SelfEff	0.058	0.425	ns
SupportIntent	Enhance	0.291	<0.001	***
SupportIntent	EffAvoid	−0.100	0.052	marginal
SubstitutionIntent	SRL	−0.168	<0.05	*
SubstitutionIntent	SelfEff	−0.153	<0.05	*
SubstitutionIntent	Enhance	−0.236	<0.01	**
SubstitutionIntent	EffAvoid	0.224	<0.01	**

Note: Std. B = standardized path coefficient. *** $p < 0.001$; ** $p < 0.01$; * $p < 0.05$; marginal: $0.05 \leq p < 0.10$; ns = $p \geq 0.10$.

As shown in Table I, Enhancement Orientation was a significant positive predictor of SRL, Self-Efficacy, and Support Intention, while also negatively predicting Substitution Intention. Efficiency-Avoidance Orientation significantly predicted higher Substitution Intention but had no significant direct effect on SRL or Self-Efficacy. Furthermore, SRL positively influenced Self-Efficacy and Support Intention while reducing Substitution Intention. Finally, Self-Efficacy significantly suppressed Substitution Intention but did not predict Support Intention. The model accounted for meaningful variance in SRL ($R^2 = 0.44$), Self-Efficacy ($R^2 = 0.26$), Support Intention ($R^2 = 0.36$), and Substitution Intention ($R^2 = 0.32$).

C. Indirect Effects and Mediation

Analysis of indirect effects confirmed the mediating role of SRL. A significant positive indirect effect was found from Enhancement Orientation to Support Intention via SRL (Standardized Indirect Effect = 0.19). A significant negative indirect effect was also found from Enhancement Orientation to Substitution Intention via SRL (Standardized Indirect Effect = −0.11). Other potential indirect paths were non-significant or negligible.

D. Robustness and Hypothesis Summary

Robustness checks affirmed the model’s stability. Variance Inflation Factors (VIFs) for all predictors ranged from 1.13 to 2.29, indicating that multicollinearity was not a concern. Procedural remedies and the excellent model fit

also suggest that common method bias did not substantially distort the findings.

The measurement model demonstrates strong reliability and discriminant validity, and the overall model fit is exemplary. Enhancement Orientation exerts broad, beneficial influence, operating directly and via SRL to elevate supportive intention and dampen substitution intention. SRL and Self-Efficacy form layered defensive pathways against substitution, while Self-Efficacy does not add incremental prediction for supportive intention. Efficiency Avoidance Orientation chiefly increases substitution intention directly and does not significantly undermine regulatory or efficacy resources. These results provide a robust empirical foundation for subsequent mechanism refinement and targeted intervention design.

V. DISCUSSION

Relative to historical surveys of frequency-centered artifacts [1, 2], the distinction of intention levels focuses on the psychological substrate underlying visible traces. When the cognitive offloading lens, as applied by Risko and Gilbert [3], is used, we recognize that scaffolded offloading (assisting) is distinct in quality from substitutional delegation. Self-efficacy's selective inhibition (as opposed to activation) profile revises social-cognitive explanations, such as those from Schunk and DiBenedetto [10]. Situating the findings in academic integrity discourse [11, 12] sidesteps the emphasis on low-hanging motivational fruit downstream in favor of higher-hanging motivational fruit upstream.

Three instructional levers are identified from a pedagogical perspective: enhancement framing to foster long-term competency; SRL micro skill scripting (e.g., focusing on process verbs and using comparison-evaluation checklists) to trigger enabling agency; and efficacy messaging to position learners as competent creators and evaluators, thereby mitigating unwarranted outsourcing. AI use disclosures can also structure metacognitive checkpoints. Intention scales could be deployed early in the diagnostic for high-risk profiles (e.g., high efficiency-avoidance with moderate SRL) to provide adaptive scaffolding.

Limitations include the cross-sectional self-report nature of the data, the absence of trace or performance data (so the translation from intention to enactment is inferential), a possible conflation of adaptive and maladaptive efficiency motives, and a domain focus on writing. Furthermore, the findings are susceptible to social desirability and miscalibration biases, as identified by Paulhus [14] and Kruger and Dunning [15], respectively.

For future work, we urge that interaction traces (prompt history, iteration number, human edit frequency) and outcome measures (quality difference, integrity threat signs) be incorporated. We also recommend longitudinal or cross-lagged designs to establish temporal precedence, decompositions of the motives for efficiency, and the validation of interventions into efficient helping (e.g., enhancement reframing and micro skill recipes). As has been called for in multimodal SRL analysis by Azevedo and Gašević [16], these extensions could also describe a

complete pipeline ranging from motivational factors to regulatory activities, goals, acted behaviors, and academic achievements.

VI. CONCLUSION

This study advances understanding of academic AI use by showing that two motivational orientations (enhancement vs. efficiency-avoidance) and two regulatory constructs (SRL, self-efficacy) jointly shape divergent intentions: supportive partnership versus substitutional outsourcing. Enhancement Orientation, operating largely through SRL (and modestly through self-efficacy), increases Supportive Intention and suppresses Substitution Intention. Efficiency-avoidance exerts a direct positive link with substitution without markedly undermining SRL. Self-efficacy functions defensively—reducing substitutional intention—yet does not significantly elevate supportive intention.

The mechanism is therefore asymmetric. An enhancement → SRL pathway converts growth motives into deliberate planning, monitoring, and evaluative reflection [3], which simultaneously encourages scaffolding uses of AI and discourages wholesale outsourcing. Efficiency-avoidance, by contrast, supplies motivational permission for direct delegation even when regulatory capacity is not fully diminished. This helps explain why generic SRL training alone may not curb over-reliance: the shortcut motive remains intact unless reframed.

Taken together, separate subgoals of collaboration vs. replacement and their absorption in an interactional control-support balance sharpen explanatory precision beyond frequency of use. Effectiveness and ineffectiveness (perceived as control and replacement) tend to preserve a collaborative bond and act to ineluctably replace, respectively; self-efficacy favors preserving control but is likely to be utilized principally to ensure self-preservation from too much outsourcing. These observations provide operational guidelines for learning design and academic integrity management, pending future validation through interaction traces.

CONFLICT OF INTEREST

The author declares no conflict of interest.

REFERENCES

- [1] E. Kasneci *et al.*, "ChatGPT for good? On opportunities and challenges of large language models for education," *Computers & Education: Artificial Intelligence*, vol. 2, 100114, 2023.
- [2] B. D. Lund, T. Wang, N. R. Mannuru, B. Nie, S. Shimray, and Z. Wang, "ChatGPT and a new academic reality: Artificial intelligence-written research papers and the ethics of authorship," *Journal of the Medical Library Association*, vol. 111, no. 3, pp. 434–441, 2023.
- [3] B. J. Zimmerman, "Attaining self-regulation: A social cognitive perspective," in *Proc. Handbook of Self-Regulation*, M. Boekaerts, P. R. Pintrich, and M. Zeidner, Eds. Academic Press, 2000, pp. 13–39.
- [4] E. F. Risko and S. J. Gilbert, "Cognitive offloading," *Trends in Cognitive Sciences*, vol. 20, no. 9, pp. 676–688, 2016.

- [5] D. Wood, J. S. Bruner, and G. Ross, "The role of tutoring in problem solving," *Journal of Child Psychology and Psychiatry*, vol. 17, no. 2, pp. 89–100, 1976. doi: 10.1111/j.1469-7610.1976.tb00381.x
- [6] C. S. Dweck, "Motivational processes affecting learning," *American Psychologist*, vol. 41, no. 10, pp. 1040–1048, 1986.
- [7] J. Biggs, D. Kember, and D. Y. P. Leung, "The revised two-factor Study Process Questionnaire: R-SPQ-2F," *British Journal of Educational Psychology*, vol. 71, no. 1, pp. 133–149, 2001. doi: 10.1348/000709901158433
- [8] A. J. Elliot and M. A. Church, "A hierarchical model of approach and avoidance achievement motivation," *Journal of Personality and Social Psychology*, vol. 72, no. 1, pp. 218–232, 1997. doi: 10.1037/0022-3514.72.1.218
- [9] C. Midgley, A. Kaplan, M. Middleton, M. L. Maehr, T. Urdan, L. H. Anderman, E. Anderman, and R. Roeser, "The development and validation of scales assessing students' achievement goal orientations," *Contemporary Educational Psychology*, vol. 23, no. 2, pp. 113–131, 1998. doi: 10.1006/ceps.1998.0965
- [10] D. H. Schunk and M. K. DiBenedetto, "Motivation and social cognitive theory," *Contemporary Educational Psychology*, vol. 60, pp. 101832, 2020. doi: 10.1016/j.cedpsych.2019.101832
- [11] T. Bretag, R. Harper, M. Burton, C. Ellis, P. Newton, P. Rozenberg, S. Saddiqui, and K. van Haeringen, "Contract cheating: A survey of Australian university staff," *Studies in Higher Education*, vol. 44, no. 11, pp. 1857–1873, 2019. doi: 10.1080/03075079.2018.1462789
- [12] D. R. E. Cotton, P. A. Cotton, and J. R. Shipway, "Chatting and cheating: Ensuring academic integrity in the era of ChatGPT," *Innovations in Education and Teaching International*, vol. 61, no. 2, pp. 228–239, 2023. doi: 10.1080/14703297.2023.2190148
- [13] P. R. Pintrich, D. A. F. Smith, T. Garcia, and W. J. McKeachie, "Reliability and predictive validity of the Motivated Strategies for Learning Questionnaire (MSLQ)," *Educational and Psychological Measurement*, vol. 53, no. 3, pp. 801–813, 1993. doi: 10.1177/0013164493053003024
- [14] D. L. Paulhus, "Measurement and control of response bias," in *Proc. Measures of Personality and Social Psychological Attitudes*, J. P. Robinson, P. R. Shaver, and L. S. Wrightsman, Eds. Academic Press, 1991, pp. 17–59.
- [15] J. Kruger and D. Dunning, "Unskilled and unaware of it: How difficulties in recognizing one's own incompetence lead to inflated self-assessments," *Journal of Personality and Social Psychology*, vol. 77, no. 6, pp. 1121–1134, 1999.
- [16] R. Azevedo and D. Gašević, "Analyzing multimodal multichannel data about self-regulated learning with advanced learning technologies: Issues and challenges," *Computers in Human Behavior*, vol. 96, pp. 207–210, 2019. doi: 10.1016/j.chb.2019.03.025

Copyright © 2026 by the authors. This is an open access article distributed under the Creative Commons Attribution License which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited ([CC BY 4.0](https://creativecommons.org/licenses/by/4.0/)).