

ARTICLE

Reciprocal Prompt Co-Adaptation with GPT-4 for
entrepreneurship educationQinjie Shen*, Salin Pituksung, Lakkamol Atsawamaitree, and Panupong Pituksung*

Department of Bachelor of Business Administration, International College, Raffles International College Bangkok, Samut Prakan, Thailand

Abstract

Generative artificial intelligence (AI) is increasingly used in education, but its role in the development of an entrepreneurial mindset remains underexplored. We propose Reciprocal Prompt Co-Adaptation (RPCA), a closed-loop human–AI scaffolding framework that supports the development of an entrepreneurial mindset through iterative prompt refinement between students and GPT-4. In this framework, “reciprocal” refers to interaction-level adaptation: students provide ratings and reflections, and the system revises prompts and coaching strategies accordingly. It does not imply that GPT-4 learns in the human sense or updates its model parameters. Unlike static or one-sided adaptive systems, RPCA integrates three components: (i) a prompt alignment module that incorporates student ratings and reflections through meta-prompting, (ii) an entrepreneurship-specific reward-shaping engine inspired by reinforcement learning from human feedback, and (iii) a confidence tracking layer that adjusts scaffolding and challenge based on textual, behavioral, and brief self-report signals. In an eight-week randomized experiment with 120 Thai undergraduates across four conditions, namely RPCA, static templates, adaptive prompting, and human–AI hybrid, RPCA produced significantly larger pre-to-post gains in opportunity recognition, risk propensity, and creative self-efficacy compared with all baseline conditions. It also achieved higher session-level prompt relevance and faster convergence to effective prompts. An embedded ablation study further suggested that removing reward shaping reduced advantages in opportunity recognition and creative self-efficacy, whereas removing confidence tracking reduced creative self-efficacy and perceived prompt quality. To support transparency and replicability, the interaction protocol, rubric-guided scoring criteria, and illustrative examples of prompt revision are specified in the appendices. These findings provide initial evidence that off-the-shelf large language models, when embedded in structured, feedback-driven tutoring protocols, can serve as adaptive, dialogic scaffolds rather than static content providers, offering a scalable approach to supporting the development of a nonlinear mindset in entrepreneurship education.

*Corresponding authors:

Qinjie Shen
(SHENQINJIE@
rafflesinternationalcollege.ac.th);
Panupong Pituksung
(panupongp@raffles.education)

Citation: Shen, Q., Pituksung, S.,
Atsawamaitree, L. & Pituksung,
P. (2026). Reciprocal Prompt
Co-Adaptation with GPT-4 for
entrepreneurship education. *Int
J Systematic Innovation*, 10(4):
026110022.

[https://doi.org/10.6977/
IJoSI.202608_10\(4\).0004](https://doi.org/10.6977/IJoSI.202608_10(4).0004)

Received: March 10, 2026

Revised: June 5, 2026

Accepted: June 22, 2026

Published online: July 8, 2026

Copyright: © 2026 Author(s).
This is an Open-Access article
distributed under the terms of the
Creative Commons Attribution
License, permitting distribution,
and reproduction in any medium,
provided the original work is
properly cited.

Publisher's Note: AccScience
Publishing remains neutral with
regard to jurisdictional claims in
published maps and institutional
affiliations.

Keywords: Entrepreneurship education; Entrepreneurial mindset; Human–artificial intelligence interaction; Large language models; Prompt engineering; Creative self-efficacy

1. Introduction

The integration of artificial intelligence (AI) into education has reshaped how learners receive feedback, practice complex skills, and engage with open-ended tasks. This shift is particularly relevant to entrepreneurship education, where learning outcomes extend beyond the acquisition of conceptual or procedural knowledge to include opportunity recognition (OR), calculated risk-taking, creative self-efficacy (CSE), and persistence under uncertainty. Although AI tutors have shown promise in supporting knowledge delivery and task-based guidance (Létourneau *et al.*, 2025), their use in entrepreneurial mindset development remains less well understood. Many existing systems rely on fixed prompts or one-sided adaptive mechanisms that primarily rely on student performance metrics (Diyab *et al.*, 2025; D. Lee & Palmer, 2025). Such designs may provide useful support, but they offer limited opportunities for students to actively shape the prompts, feedback, and scaffolding that guide their learning.

Entrepreneurship education requires iterative experimentation, reflection, and revision. Students are often expected to generate ideas, test assumptions, respond to feedback, and refine their reasoning in uncertain contexts. Recent studies have noted that conventional intelligent tutoring systems are less well suited to these higher-order and dispositional learning goals because they often prioritize accuracy, efficiency, or content delivery rather than learner agency and reflective adaptation (Chen *et al.*, 2024; Roll & Wylie, 2016). In this context, generative AI may be useful not because it learns in the same way as a human peer, but because its prompts and feedback strategies can be repeatedly adjusted in response to learner input. This makes it possible to design tutoring protocols in which students are not merely recipients of AI-generated content, but active contributors to the structure and direction of the interaction.

To address this design problem, we introduce Reciprocal Prompt Co-Adaptation (RPCA), a closed-loop human-AI scaffolding framework for entrepreneurship education. In this study, “reciprocal” refers to interaction-level adaptation rather than mutual learning between the student and the AI system. GPT-4 is not assumed to acquire knowledge, update its model parameters, or demonstrate learning in the human sense. Instead, RPCA operationalizes reciprocity as an iterative process in which students provide ratings and reflections, and the system uses this feedback to revise prompts and adjust coaching strategies. The framework is therefore positioned as a form of AI-supported dialogic scaffolding and prompt co-adaptation, rather than as evidence of AI-side learning.

Reciprocal Prompt Co-Adaptation contains three main components. First, the prompt alignment module establishes a structured protocol for iterative prompt refinement. Students evaluate the relevance of prompts and provide short reflections, which are then incorporated through meta-prompting to revise subsequent prompts. Second, an entrepreneurship-specific reward-shaping engine, inspired by reinforcement learning from human feedback (RLHF) (Chaudhari *et al.*, 2025; González Barman *et al.*, 2025), guides feedback toward entrepreneurial mindset components such as divergent thinking and calculated risk-taking, rather than optimizing only for correctness or engagement. Third, a confidence tracking layer adjusts the level of scaffolding and challenge based on textual, behavioral, and brief self-report signals of the learner’s state. Together, these components create a feedback-driven tutoring protocol that keeps the interaction aligned with both pedagogical goals and student responses.

The proposed framework extends prior work in human-AI interaction and adaptive tutoring in several ways. Unlike static prompt templates, RPCA allows prompts to evolve over the course of interactions in response to student feedback. Unlike one-sided adaptive prompting, it gives learners a structured role in shaping the direction of AI-supported tutoring. Unlike generic AI tutoring systems, it incorporates entrepreneurship-specific reward criteria that target OR, creative exploration, and calculated risk-taking. Finally, its confidence-aware adaptation layer enables the system to adjust scaffolding intensity when students appear uncertain or ready for greater challenge. These features are intended to support learner agency and the development of a nonlinear mindset while remaining deployable with an off-the-shelf large language model (LLM) and without fine-tuning model parameters.

We evaluate RPCA in an eight-week classroom experiment in a Business Creativity and Innovation course with Thai undergraduate students. The study compares RPCA-enabled GPT-4 tutoring with three baseline conditions: static templates, adaptive prompting, and human-AI hybrid prompting. The evaluation focuses on both learning outcomes and the quality of the interaction process, including pre-to-post changes in OR, risk propensity (RP), and CSE, as well as session-level prompt relevance and convergence speed. In addition, an embedded ablation study examines the contributions of the reward-shaping engine and the confidence tracking layer. By doing so, this study provides initial evidence on whether feedback-driven prompt co-adaptation can more effectively support the development of an entrepreneurial mindset than static or one-sided AI-assisted prompting approaches.

Guided by the goal of developing and evaluating a deployable human–AI prompt co-adaptation protocol for entrepreneurship education, this study addresses three research questions (RQs):

- (i) RQ1: Does RPCA lead to larger pre-to-post gains in entrepreneurial mindset outcomes, namely OR, RP, and CSE, than baseline AI-assisted prompting conditions?
 - H1–H3: RPCA will yield larger gains than static templates, adaptive prompting, and human–AI Hybrid prompting on OR (H1), RP (H2), and CSE (H3).
- (ii) RQ2: Does RPCA improve interaction process quality relative to baseline prompting approaches?
 - H4: RPCA will achieve higher session-level prompt relevance and faster convergence to effective prompts than the baseline conditions.
- (iii) RQ3: Which components of RPCA contribute to the observed effects?
 - H5: Removing the reward-shaping engine will reduce OR and CSE task performance relative to full RPCA.
 - H6: Removing the confidence tracking layer will reduce CSE and perceived prompt quality, as indicated by prompt relevance, relative to full RPCA.

2. Literature review

The development of AI systems for educational purposes has increasingly moved beyond simple content delivery toward adaptive and interactive forms of learning support. This shift is particularly relevant to domains such as entrepreneurship education, where learning involves iterative idea generation, feedback interpretation, risk evaluation, and reflective revision. However, the growing use of generative AI in education also requires conceptual caution. Although LLMs can respond to learner input and generate context-sensitive feedback, this does not mean they learn in the same sense as human students or acquire new knowledge during a classroom interaction. In the present study, AI-supported adaptation is therefore understood as interaction-level prompt and scaffolding adjustment, not as model-side learning.

This study develops the conceptual foundation for RPCA in three steps. First, it reviews work on feedback-driven human–AI adaptation and learner agency in AI-supported education. Second, it discusses RLHF and adaptive prompt engineering as technical and pedagogical foundations for structured prompt refinement. Third, it connects confidence-aware scaffolding to entrepreneurial

mindset development and CSE. Together, these areas provide the basis for designing RPCA as a closed-loop prompt-adaptation framework that supports student learning through structured interaction with GPT-4 without requiring model fine-tuning.

2.1. Feedback-driven human–artificial intelligence adaptation in education

Contemporary AI-supported learning research has increasingly emphasized adaptive interaction rather than one-way content transmission. Earlier intelligent tutoring systems often focused on diagnosing knowledge gaps and delivering targeted instructional content. More recent work on hybrid human–AI learning technologies highlights the importance of co-regulation, in which learners, teachers, and AI tools contribute differently to regulating learning processes (Molenaar, 2022). In this view, learner agency remains central: AI systems can provide prompts, feedback, and analytics, but students must still interpret, evaluate, and act upon such support.

This distinction is important for the present study. RPCA does not claim that GPT-4 becomes a human-like learning partner or that it participates in reciprocal knowledge construction in the same way as a human peer. Instead, it treats GPT-4 as a configurable dialogic scaffold whose prompts and coaching strategies can be adjusted in response to student feedback. Such a design aligns with broader arguments that AI in education should augment human judgment rather than replace it (Holmes & Tuomi, 2022; Kim *et al.*, 2022). In entrepreneurship education, this distinction is especially important because students must learn to evaluate uncertainty, justify risk-taking, and revise ideas through reflection rather than simply receive correct answers.

Recent studies also suggest that conventional intelligent tutoring systems are often less suited to higher-order and dispositional learning goals, such as creativity, agency, resilience, and OR (Chen *et al.*, 2024; Roll & Wylie, 2016). These goals require more than adaptive delivery of content. They require interaction designs that allow learners to shape the direction of feedback, clarify their needs, and reflect on how prompts influence their thinking. RPCA builds on this direction by giving students a structured role in rating prompts, providing short reflections, and influencing the subsequent configuration of AI-supported tutoring.

2.2. Reinforcement learning from human feedback and educational reward design

Reinforcement learning from human feedback has become an influential approach for aligning AI system behavior

with human preferences and task goals. In educational contexts, RLHF-inspired approaches are relevant because they provide a way to formalize what kinds of system responses should be encouraged. However, educational adaptation cannot rely only on generic preferences such as helpfulness, engagement, or fluency. For complex learning domains, reward criteria need to reflect the specific learning goals of the activity.

In entrepreneurship education, this presents a particular challenge. Entrepreneurial learning involves multiple, sometimes competing, goals, including creativity, feasibility, calculated risk-taking, and responsiveness to feedback. A prompt that maximizes comfort or immediate satisfaction may not necessarily promote deeper entrepreneurial learning. Similarly, a prompt that produces fluent ideas may not encourage students to test assumptions, tolerate uncertainty, or revise weak ideas. For this reason, RPCA uses an entrepreneurship-specific reward-shaping mechanism inspired by RLHF principles (Chaudhari *et al.*, 2025; González Barman *et al.*, 2025), but it neither trains a new policy model nor updates GPT-4 parameters.

Instead, reward shaping in RPCA functions as a pedagogical evaluation layer. Student responses are evaluated using rubric-guided criteria related to divergent thinking and calculated risk-taking. These scores are then used to guide formative feedback and subsequent prompt refinement. This approach allows the system to remain deployable with an off-the-shelf LLM while making the pedagogical goals of adaptation more explicit. At the same time, because the reward scores are model-assisted and rubric-dependent, they should be interpreted as provisional indicators rather than external ground-truth measures of entrepreneurial ability. This limitation is addressed further in Sections 4 and 5.

2.3. Adaptive prompt engineering and learner agency

Prompt engineering has evolved from static template construction toward more adaptive forms of interaction design. In educational settings, prompts do not merely instruct an AI system; they shape the kinds of questions students receive, the types of reasoning they are invited to practice, and the level of challenge or scaffolding embedded in the learning activity. Human-in-the-loop approaches show that iterative feedback can refine AI outputs while keeping adaptation accountable to human judgment (Mosqueira-Rey *et al.*, 2023).

Existing work on student–AI collaboration suggests that teachers often act as designers of interaction patterns, curating prompts and workflows that help students use AI

more productively (Kim *et al.*, 2022). However, many such designs maintain a clear separation between the prompt designer and the learner. The teacher or instructional designer prepares prompts, while students respond to them. This arrangement can be useful, but it gives students limited opportunities to reflect on whether a prompt fits their learning needs or to participate in refining the interaction itself.

Reciprocal Prompt Co-Adaptation addresses this gap by treating prompt refinement as part of the learning process. Students evaluate the relevance of prompts, provide brief feedback, and influence how subsequent prompts are revised. This does not mean that the student and GPT-4 co-learn in a strong reciprocal sense. Rather, it means that students contribute feedback that changes the structure of the tutoring interaction. The pedagogical value lies in making prompt adaptation visible and participatory, thereby supporting learner agency, metacognitive awareness, and more intentional engagement with AI-generated guidance.

2.4. Confidence-aware scaffolding and mindset development

Learner confidence is a key factor in adaptive educational support. In open-ended learning contexts, students' confidence can influence whether they persist, revise their ideas, ask for help, or abandon uncertain possibilities too early. Multimodal approaches that combine behavioral traces, interaction logs, and self-reports have shown promise for capturing aspects of learner state, including engagement, confusion, and self-efficacy (Guerrero-Sosa *et al.*, 2025; Holmes & Tuomi, 2022). However, many existing applications focus on knowledge acquisition, such as STEM problem-solving or language learning, rather than on dispositional growth in entrepreneurial education.

Entrepreneurial mindset development requires careful calibration of challenge and support. If AI tutoring provides too much scaffolding, students may become dependent on examples or safe suggestions. If it provides too little support, students with lower confidence may disengage or avoid risk-taking. A confidence-aware system can help manage this tension by increasing guidance when students appear uncertain and introducing more open-ended challenges when students appear ready for greater autonomy.

In RPCA, confidence tracking is used for scaffolding decisions rather than for high-stakes assessment. The confidence tracking layer integrates textual, behavioral, and brief self-report signals to estimate whether the student may need more structure or greater challenge. This estimate does not represent a validated psychological diagnosis

of confidence. Rather, it functions as a transparent, classroom-deployable indicator for adapting the tutoring style. Section 4, therefore, reports how confidence-related features were extracted, how signals were scaled, and how thresholds were used to determine scaffolding responses.

2.5. Entrepreneurial mindset, creative confidence, and artificial intelligence-supported scaffolding

Developing an entrepreneurial mindset involves more than acquiring business knowledge. It requires cultivating a cognitive and behavioral orientation toward OR, value creation, uncertainty, and iterative action. Empirical work highlights characteristics such as tolerance for ambiguity, proactive problem solving, and the ability to pivot strategies based on feedback (A. Lee & Jung, 2021; Sandhu, Sarkar, & Subburaj, 2025). These traits are not fixed personality attributes; they can be strengthened through appropriately designed educational experiences.

Entrepreneurial dispositions often develop through cycles of action and reflection embedded in authentic tasks. Students generate ideas, receive feedback, test assumptions, and revise their reasoning. This process is path-dependent because learners interpret successes and failures in light of prior experience, perceived control, and social context. As a result, entrepreneurial mindset development does not always follow a linear trajectory. Students may need to revisit earlier ideas, reconsider assumptions, and recombine partial insights before making meaningful progress.

This nonlinear developmental pattern creates challenges for AI-assisted learning systems. Many AI tutors are designed to support relatively structured learning progressions, where the system can identify correct responses, provide explanations, and move students toward predefined mastery. However, entrepreneurial learning often involves ambiguity and multiple acceptable pathways. Recent reviews of AI-driven tutoring indicate that many systems remain optimized for short-term task performance, with less attention to longer-term dispositional growth (Létourneau *et al.*, 2025). RPCA addresses this challenge by combining prompt refinement, domain-specific reward shaping, and confidence-aware scaffolding into a single interaction protocol.

2.6. Creative self-efficacy and calibrated challenge

Creative self-efficacy refers to learners' beliefs in their ability to generate novel and valuable ideas. It is particularly important in entrepreneurship education because students must not only produce ideas but also believe that they can improve them through feedback and experimentation. Contemporary creativity research conceptualizes creative behavior as agentic action, in which individuals decide

to enact their creative potential based on self-beliefs and perceived opportunities (Karwowski & Beghetto, 2019). Related work on creative potential suggests that confidence beliefs help translate creative capacity into observable creative action (Richard, Holder, & Cairney, 2021).

Students with stronger CSE are more likely to persist in ambiguous problem-solving situations and to tolerate failure as part of the entrepreneurial process. In contrast, students with weaker creative confidence may abandon ideas prematurely or choose safe options even when they possess relevant knowledge. Studies on creative mindsets further suggest that growth-oriented beliefs about creativity can support persistence in the face of difficulty and help students interpret setbacks as opportunities for improvement (He & Chiang, 2024). Conceptual work on creative flourishing similarly emphasizes the need to balance challenge, self-efficacy, and environmental support (Glass, 2025).

These insights create a central design problem for AI-supported entrepreneurship education: how can the system provide enough scaffolding to prevent discouragement while still preserving productive uncertainty and challenge? RLHF pipelines optimized mainly for positive ratings or user satisfaction may risk favoring comfort over growth (Chaudhari *et al.*, 2025). RPCA addresses this problem by using confidence-aware adaptation to calibrate scaffolding intensity and reward shaping, keeping feedback aligned with entrepreneurial learning goals.

2.7. Implications for the Reciprocal Prompt Co-Adaptation framework

The literature reviewed above suggests three design requirements for AI-supported entrepreneurial mindset development. First, the system should accommodate nonlinear learning trajectories by allowing learners to revise, revisit, and recombine ideas over time. Second, it should balance challenge and support so that CSE can develop through productive struggle rather than simple reassurance. Third, it should give learners a structured role in shaping AI-supported tutoring, so that adaptation is not only generated by the system but also informed by student reflection.

Reciprocal Prompt Co-Adaptation operationalizes these requirements through a closed-loop prompt-adaptation framework. The prompt alignment module allows students to influence prompt revision through ratings and reflections. The reward-shaping engine keeps feedback aligned with entrepreneurship-specific criteria such as divergent thinking and calculated risk-taking. The confidence tracking layer adjusts the level of scaffolding

and challenge based on learner-state signals. Importantly, these mechanisms do not imply that GPT-4 learns in the human sense. The adaptive element lies in the interaction protocol, prompt revision process, and coaching strategy, while student learning remains the primary empirical outcome of interest.

3. The Reciprocal Prompt Co-Adaptation framework

The RPCA framework establishes a structured interaction protocol between students and GPT-4 through three interconnected components: the prompt alignment module, the entrepreneurship-specific reward-shaping engine, and the confidence tracking layer. These components form a closed-loop system in which prompt refinement is guided by student feedback, domain-specific pedagogical criteria, and confidence-aware scaffolding decisions. In our implementation, GPT-4 serves as the backbone model, while RPCA governs prompt revision, entrepreneurial reward scoring, and coaching strategy adjustment. The framework is deployable via standard GPT-4 access and requires no fine-tuning or parameter updates.

Let u_t denote the student's input at interaction turn t , y_t denote the AI tutor's response, and P_t denote the active system prompt that configures GPT-4's behavior at that turn. Student feedback is represented as $F_t = (q_t, s_t)$, where q_t is the prompt-relevance rating and s_t is the student's free-text suggestion. The entrepreneurship-specific composite reward is denoted as $Reward_t$, and the confidence estimate is denoted as C_t . RPCA maintains a sequence of prompts $\{P_t\}$ that can be revised over the course of interactions based on student ratings, student reflections, and pedagogical criteria. The purpose of this revision process is not to train GPT-4, but to adapt the instructional conditions under which GPT-4 generates subsequent responses.

3.1. Prompt alignment module

The prompt alignment module formalizes student participation in prompt refinement. Each interaction begins with an active prompt P_t designed to support entrepreneurial thinking. After engaging with GPT-4 on this prompt, students provide structured feedback in the form of $F_t = (q_t, s_t)$, where q_t represents a 5-point rating of the prompt's relevance to the learning task, and s_t contains a short free-text reflection or suggestion for improvement. Examples of such suggestions include "give more concrete examples," "ask more challenging questions," or "help me compare two ideas."

The system processes this feedback through two pathways. First, the relevance rating q_t indicates whether

the current prompt should be retained, slightly modified, or substantially revised. Second, the reflection s_t provides qualitative information about what kind of adjustment may be needed, such as more scaffolding, more examples, more constraints, or more open-ended challenges. For expository clarity, semantic alignment between the student's reflection and the current prompt can be represented using cosine similarity (Equation 1):

$$\text{Sim}(s_t, P_t) = \cos(e_s, e_p) \quad (1)$$

where e_s and e_p denote sentence embeddings of the student reflection and the current prompt. This similarity score is not treated as evidence of AI learning; rather, it provides a transparent indicator of how closely the student's suggested direction aligns with the existing prompt.

The updated prompt can be represented as Equation 2:

$$P_{t+1} = \text{GPT-4 edit}(P_t, \alpha q_t + [1 - \alpha]\text{Sim}(s_t, P_t)) \quad (2)$$

where α controls the relative contribution of the structured relevance rating and the embedding-based semantic alignment score. In the present implementation, this update rule is used to guide meta-prompt editing rather than to train a model. The exact setting of α and the rules for prompt revision are reported in Section 4 to support transparency and replicability.

3.2. Entrepreneurship-specific reward-shaping engine

The reward-shaping engine introduces domain-specific feedback criteria into the RPCA framework. Rather than optimizing only for correctness, fluency, or engagement, RPCA evaluates student responses in relation to entrepreneurial mindset goals. Specifically, GPT-4 is used as a rubric-guided evaluator to score student responses on two dimensions: creativity and calculated risk-taking. This design draws on multidimensional reward-shaping principles from RLHF research (Chaudhari *et al.*, 2025; Qian, 2025), but it does not involve training a new reinforcement-learning policy.

Given a student input u_t , the system produces two rubric-aligned scores: (i) $\text{CreativityScore}(u_t)$: the novelty, diversity, and elaboration of ideas, and (ii) $\text{RiskScore}(u_t)$: the student's willingness to consider unconventional but justifiable options, including awareness of uncertainty, trade-offs, and feasibility. Both scores are normalized to a 0–100 scale. RPCA then combines them into a composite reward (Equation 3):

$$\text{Reward}_t = \beta \text{CreativityScore}(u_t) + (1 - \beta) \text{RiskScore}(u_t) \quad (3)$$

where β controls the relative emphasis on divergent thinking and calculated risk-taking.

The composite reward Reward_t serves three pedagogical functions. First, it supports formative feedback to students by identifying whether their responses show creative exploration, risk awareness, or both. Second, it provides instructors with interpretable process indicators about class-level entrepreneurial thinking. Third, it informs subsequent prompt refinement by identifying which prompt patterns tend to elicit responses aligned with the learning goals. Because these scores are generated through a rubric-guided, model-assisted process, they should be interpreted with caution and supplemented with external or human evaluation where possible.

3.3. Confidence tracking layer

The confidence tracking layer estimates the level of scaffolding a student may need during the interaction. It is designed as a transparent classroom-deployable estimator rather than a fully trained psychometric model. The estimate is used only to adjust tutoring style, not to make high-stakes judgments about students.

3.4. Combination of three signals in Reciprocal Prompt Co-Adaptation

TextSignal_t are indicators extracted from student responses, such as explicit confidence statements, uncertainty markers, help-seeking language, willingness to revise, and evaluative certainty. These features are coded using a fixed rubric implemented through GPT-4.

BehaviorSignal_t are interaction-log features such as response latency, the number of voluntary follow-up turns, and the number of revision attempts. These signals are normalized within a session to reduce the influence of device or network differences.

SelfReport_t are brief momentary confidence ratings, such as “How confident do you currently feel about generating valuable business ideas?” rated on a 1 to 5 scale and rescaled to a 0 to 1 range.

The confidence estimate can be represented as Equation 4:

$$C_t = \sigma(w_1 \text{TextSignal}_t + w_2 \text{BehaviorSignal}_t + w_3 \text{SelfReport}_t) \quad (4)$$

where σ maps the value into the range $[0, 1]$, and w_1 , w_2 , and w_3 determine the contribution of textual, behavioral, and self-report signals. These weights are not learned through supervised model training in the present prototype. Instead, they are specified in advance for pedagogical interpretability and reported in Section 4. This choice

keeps the system feasible for classroom deployment while making the limitations of the confidence estimate explicit.

The resulting C_t value determines the coaching strategy. Lower confidence estimates trigger more scaffolding, concrete examples, and step-by-step guidance. Moderate confidence estimates lead to balanced guidance that combines support with reflective questioning. Higher confidence estimates trigger more open-ended prompts, greater challenge, and stronger invitations to explore bolder entrepreneurial possibilities. This adaptive scaffolding mechanism is intended to support CSE by balancing reassurance with productive challenge.

3.5. GPT-4 meta-prompt editing without fine-tuning

The GPT-4 edit component implements structured prompt revision without modifying GPT-4's model weights. Instead of fine-tuning, RPCA uses a meta-prompt editing approach in which GPT-4 is instructed to revise the active prompt P_t according to constraints derived from student feedback, pedagogical objectives, and reward-shaping criteria. This approach avoids the infrastructure requirements of parameter-efficient fine-tuning methods such as LoRA (Hu *et al.*, 2022) and enables deployment via standard GPT-4 interfaces.

At each refinement step, GPT-4 receives three inputs: (i) the current system prompt P_t ; (ii) the combined feedback signal based on student rating and semantic alignment; and (iii) an editing rubric specifying what should be preserved, strengthened, simplified, or made more challenging.

GPT-4 then produces a revised prompt, P_{t+1} , which is checked against the editing rubric. The process can be represented as Equation 5:

$$P_{t+1} = \text{GPT-4 edit}(P_t, \alpha q_t + [1 - \alpha] \text{Sim}(s_t, P_t)) \quad (5)$$

This meta-prompt strategy improves transparency because changes to the prompt can be inspected by researchers or instructors. It also clarifies the scope of adaptation: RPCA adapts prompts and scaffolding strategies, not the underlying GPT-4 model.

3.6. Closed-loop prompt co-adaptation architecture

Reciprocal Prompt Co-Adaptation integrates prompt alignment, reward shaping, and confidence tracking into a closed-loop prompt co-adaptation architecture. The system maintains two interacting feedback loops: a prompt-quality loop and a coaching-adaptation loop.

3.6.1. Loop 1: Prompt quality optimization

Student feedback $F_t = (q_t, s_t)$ enters the prompt alignment module. The relevance rating q_t and semantic alignment

score $\text{Sim}(s_t, P_t)$ guide the revision of the active prompt. This loop supports three functions: (i) maintaining pedagogical relevance; (ii) preserving coherence across prompt revisions; and (iii) allowing students to influence how the tutoring interaction is framed. This loop should be understood as an interaction-level prompt adjustment rather than mutual learning between the student and GPT-4.

3.6.2. Loop 2: Coaching adaptation via confidence tracking

The confidence tracking layer produces a scalar estimate C_t based on textual, behavioral, and self-report signals. This estimate determines whether GPT-4 should provide more scaffolding, balanced guidance, or more open-ended challenge. The goal is to adjust the level of support to the learner's current state while keeping the task aligned with entrepreneurial learning goals.

3.6.3. Loop interaction

The two loops interact during the tutoring process. Changes in prompt quality may influence student engagement and confidence. Confidence-aware coaching may influence the type of feedback students provide. Reward-shaped feedback may reinforce prompt patterns that are more likely to support creative exploration and calculated risk-taking. These interactions create a co-adaptive tutoring environment in which the system's prompts and coaching strategies change in response to student input. However, the empirical claim remains student-centered: RPCA is evaluated by examining student learning outcomes and interaction quality, not by claiming that GPT-4 itself learns. Figure 1 illustrates the closed-loop RPCA architecture, including prompt refinement, reward shaping, and

confidence-aware scaffolding.

4. Methods

To evaluate the effectiveness of the RPCA framework, we conducted an eight-week classroom experiment comparing RPCA-enabled GPT-4 tutoring with three baseline prompting conditions. The evaluation focused on two levels of evidence. First, we examined student learning outcomes related to entrepreneurial mindset development, including OR, RP, and CSE. Second, we examined the quality of the interaction process, including prompt relevance and convergence speed. The overall design is consistent with recent work emphasizing the role of OR and self-efficacy in entrepreneurship education (Hou *et al.*, 2022; López-Muñoz *et al.*, 2023).

Consistent with the RQs, the study was structured to provide quantitative evidence for RQs 1, 2, and 3, as well as hypotheses H1–H6. To address RQ1 and H1–H3, validated scales of OR, RP, and CSE were administered before and after the intervention. To address RQ2 and H4, session-level prompt relevance ratings and convergence speed were logged as indicators of the quality of the interaction process. To address RQ3 and H5–H6, a within-subject ablation study was embedded in the RPCA condition during the final two weeks of the intervention. This ablation study examined whether removing the reward-shaping engine or the confidence tracking layer reduced the effectiveness of the full RPCA configuration.

4.1. Participants and context

A between-subjects experiment was conducted with four conditions. Participants were 120 undergraduate students, including 62 female students and 58 male students, enrolled

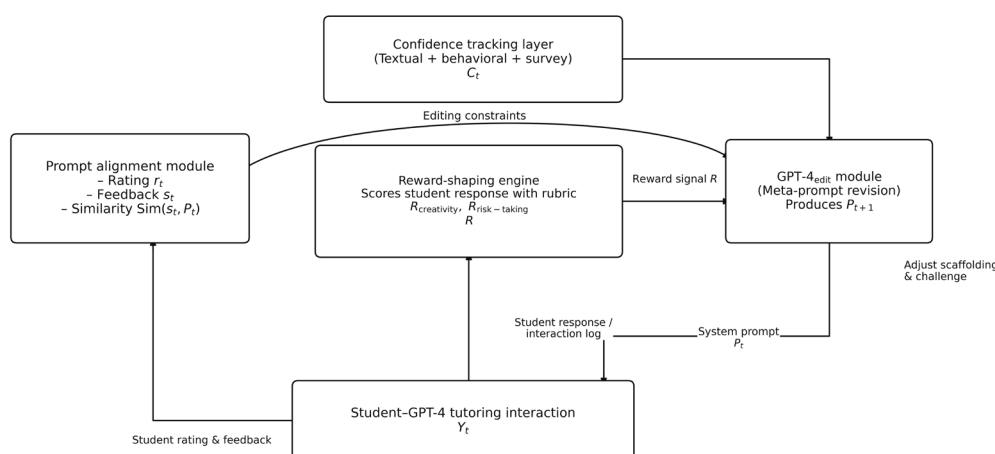


Figure 1. Reciprocal Prompt Co-Adaptation closed-loop architecture integrating prompt refinement, reward shaping, and confidence-aware scaffolding

in Business Creativity and Innovation courses at three public universities in Thailand during the 2025 academic year. The mean age was 20.8 years, with a standard deviation (SD) of 1.31. Most participants were third- or fourth-year business majors, accounting for 89% of the sample. In addition, 23% had taken prior entrepreneurship courses, and 11% reported having launched a small venture. Analyses were conducted on the full randomized sample ($N = 120$).

Within each class, students were randomly assigned to one of four conditions ($n = 30$ per condition) using a computer-generated sequence in Qualtrics, stratified by gender and prior entrepreneurship experience. To reduce cross-condition contamination, all eight weekly workshops were conducted individually on personal laptops in a dedicated computer laboratory. Students used headphones and were seated with at least one empty seat between them. Instructors received a two-hour standardized training session and followed a scripted protocol for all conditions. The four experimental conditions were as follows.

- (i) RPCA (full). Students used the full RPCA framework, including iterative prompt refinement, entrepreneurship-specific reward shaping, and confidence-aware scaffolding.
- (ii) Static templates. Students selected from a fixed menu of 25 pre-validated creativity and entrepreneurship prompts. Prompts did not change across sessions. No meta-prompt editing, reward shaping, or confidence tracking occurred.
- (iii) Adaptive prompting. GPT-4 adjusted task difficulty using a rule-based heuristic. If the previous task was completed in four or fewer turns and the student's momentary confidence was rated 4 or above on a 5-point scale, one constraint was removed to increase challenge. Otherwise, additional scaffolding was provided, such as an example idea. This condition used one-sided adaptation and did not include student-driven prompt refinement, entrepreneurship-specific reward shaping, or confidence tracking.
- (iv) Human-AI hybrid. Students alternated weekly between instructor-written prompts and unguided access to GPT-4. Prompts could be manually adjusted by the instructor, but the RPCA mechanisms were not used.

All participants, regardless of condition, engaged in eight weekly 45-minute one-on-one sessions with a GPT-4-based tutor. The conditions differed only in how prompts were selected, refined, and adapted during the tutoring process.

4.2. Implementation of Reciprocal Prompt Co-Adaptation

The RPCA condition consisted of three interconnected components: the Prompt alignment module, the entrepreneurship-specific reward-shaping engine, and the confidence tracking layer. The system was implemented using standard GPT-4 access and did not involve model fine-tuning or parameter updates. All adaptive behavior occurred through prompt revision, rubric-guided scoring, and rule-based scaffolding decisions.

4.2.1. Prompt alignment module

The prompt alignment module allowed students to participate in prompt refinement during each session. After responding to an active prompt, students rated prompt relevance on a 5-point scale and provided a short free-text suggestion. Student feedback was represented as $F_t = (q_t, s_t)$, where q_t refers to the prompt-relevance rating, and s_t refers to the student's free-text suggestion.

Prompt revision followed three rules. Ratings of 1 or 2 triggered substantial revision of the prompt. A rating of 3 triggered moderate revision. Ratings of 4 or 5 led to minor refinement or retention of the existing prompt. The free-text suggestion was used to identify the direction of revision, such as adding examples, narrowing the task, increasing the challenge, or making the prompt more concrete.

The normalized prompt-relevance rating and the semantic alignment score were weighted equally in the prompt revision process, with $\alpha = 0.50$. The 5-point prompt-relevance rating q_t was first normalized to the range 0 to 1. To estimate semantic alignment, the student's free-text suggestions s_t and the active prompt P_t were encoded into sentence-level embedding vectors using Sentence-BERT (Reimers & Gurevych, 2019). Cosine similarity was then calculated between the two vectors and used as $\text{Sim}(s_t, P_t)$. Because cosine similarity can range from -1 to 1 , the similarity score was linearly rescaled to the 0 to 1 range before being combined with the normalized prompt-relevance rating. This score represented the degree to which the student's suggested revision direction was semantically close to the active prompt.

Prompt revision was guided by the combined signal $\alpha q_t + (1 - \alpha)\text{Sim}(s_t, P_t)$. In the present implementation, α was set at 0.50 to give equal priority to structured student ratings and semantic alignment. When the combined signal was low, GPT-4 edit was instructed to substantially revise the active prompt by changing its framing, level

of scaffolding, or task constraints. When the combined signal was moderate, GPT-4 edit produced a moderate refinement, such as adding examples, narrowing the task, or increasing concreteness. When the combined signal was high, the prompt was retained or only lightly refined. The free-text suggestion was also provided directly to GPT-4 edit to ensure the revision direction remained interpretable and responsive to student feedback.

This embedding-based alignment procedure was used as a prompt-revision support mechanism rather than as evidence of AI learning. It did not fine-tune GPT-4, update model parameters, or train a new adaptive model. Instead, it provided a transparent process indicator for deciding how strongly the active prompt should be revised across interaction turns.

4.2.2. Entrepreneurship-specific reward-shaping engine

The reward-shaping engine provided domain-specific feedback criteria for evaluating student responses. Rather than optimizing for general fluency, correctness, or user satisfaction, the reward function focused on indicators of entrepreneurial mindset. GPT-4 scored student responses using a fixed rubric on two 0 to 100 dimensions: CreativityScore and RiskScore.

CreativityScore was calculated as the average of four equally weighted dimensions: novelty, diversity of ideas, elaboration, and feasibility–innovation balance. Novelty referred to the originality of the proposed idea. Diversity captured whether the student explored multiple possibilities rather than a single obvious solution. Elaboration referred to the level of detail and development in the response. Feasibility–innovation balance assessed whether the idea combined creative potential with practical plausibility.

RiskScore was also calculated as the average of four dimensions: calculated risk, uncertainty awareness, trade-off reasoning, and avoidance of reckless risk. Calculated risk captured whether the student was willing to consider uncertain but justifiable options. Uncertainty awareness was assessed by determining whether the student recognized relevant market, user, or implementation uncertainties. Trade-off reasoning evaluated whether the student explained the benefits and costs of different choices. Avoidance of reckless risk ensured that high-risk ideas were not rewarded when they lacked justification.

The composite reward was calculated as follows (Equation 6):

$$\text{Reward}_t = 0.55 \times \text{CreativityScore}(u_t) + 0.45 \times \text{RiskScore}(u_t) \quad (6)$$

In this formula, u_t denotes the student's response at interaction turn t . The slightly higher weight for creativity reflected the course's emphasis on idea generation and business creativity. Reward_t was first produced on a 0 to 100 scale and then normalized to a 0 to 1 scale for process analysis and convergence-speed calculation.

To improve scoring reliability, the research team manually reviewed a random sample of 15% of student responses and compared the rubric-based GPT-4 scores with human judgments. This check was used to identify inconsistent scoring patterns and to ensure that the scoring rubric was applied as intended. The model-generated scores were treated as process indicators rather than external ground-truth assessments of entrepreneurial ability.

4.2.3. Confidence tracking layer

The confidence tracking layer was used to adjust the level of scaffolding and challenge during tutoring. It was designed as a transparent classroom-deployable estimator rather than a trained machine-learning model or validated psychometric instrument. The estimate was used only to guide tutoring style and was not used for grading or high-stakes evaluation.

The confidence estimate combined three types of signals: TextSignal_t , BehaviorSignal_t , and SelfReport_t .

TextSignal_t was extracted from each student response using GPT-4 with a fixed scoring rubric. The rubric identified five textual indicators: uncertainty markers, help-seeking or self-doubt, self-confidence statements, hesitation or vague language, and revision willingness. Uncertainty markers included expressions such as “maybe,” “not sure,” “I think,” “perhaps,” and “kind of.” Help-seeking or self-doubt included requests for help, expressions of difficulty, or negative self-evaluation. Self-confidence statements included explicit positive expressions such as “I’m confident,” “this is solid,” or “this is a strong idea.” Hesitation or vague language included filler expressions, repetition, and overly broad descriptions. Revision willingness referred to the student's openness to iteration, improvement, and feedback. GPT-4 assigned continuous scores from 0 to 1 using the fixed rubric, and the final TextSignal_t score was calculated as a weighted and normalized composite. This procedure was rubric-guided rather than based on open-ended generation or simple keyword matching.

BehaviorSignal_t was derived from interaction logs. It included standardized response latency, voluntary follow-up turns, revision attempts, and task completion efficiency. Response latency was standardized within a session to reduce the influence of network or device differences. Voluntary follow-up turns were captured

to indicate whether students actively continued the conversation beyond the minimum task requirements. Revision attempts captured how often students revised or improved their ideas. Task completion efficiency was calculated as the number of turns required to complete the task and treated as a negative indicator when excessive turns indicated difficulty rather than productive exploration.

SelfReport_t was collected at the end of each session through brief momentary confidence items. The central item was: “How confident do you currently feel about generating valuable business ideas?” Students responded on a 1–5 Likert scale. Additional brief items captured perceived task manageability and perceived readiness to continue developing the idea. Self-report scores were normalized to the 0–1 range.

The confidence estimate was calculated as follows (Equation 7):

$$C_t = 0.35 \times \text{TextSignal}_t + 0.25 \times \text{BehaviorSignal}_t + 0.40 \times \text{SelfReport}_t \quad (7)$$

The weights were specified a priori for pedagogical interpretability rather than learned from supervised training. Self-report received the highest weight because it directly captured the student’s perceived confidence, while textual and behavioral signals provided complementary evidence from interaction data.

The resulting C_t value determined the scaffolding mode. When $C_t < 0.45$, the system entered low-confidence mode and provided stronger scaffolding, concrete examples, and step-by-step guidance. When $0.45 \leq C_t \leq 0.70$, the system entered moderate-confidence mode and provided balanced guidance that combined support with reflective questioning. When $C_t > 0.70$, the system entered high-confidence mode and increased challenge by using more open-ended prompts, inviting bolder experimentation, or asking students to test assumptions under more uncertain conditions.

Textual confidence indicators were extracted using anonymized student texts. The research team randomly reviewed 20% of the textual confidence outputs to check whether the rubric was applied consistently. This manual check was used as a quality-control procedure rather than as a separate validation study.

4.3. Measures and instruments

All scales demonstrated good to excellent internal consistency in the current sample, as shown in Table 1. These measures correspond directly to the research questions and hypotheses. OR, RP, and CSE addressed RQ1 and H1–H3. Prompt relevance and convergence speed

addressed RQ2 and H4. The ablation outcomes addressed RQ3 and H5–H6.

Table 1. Measurement instruments and reliability

Construct	Source	Items	α pre-test	α post-test	Raw scale
Opportunity recognition	Kuckertz <i>et al.</i> (2017)	5	0.89	0.91	1–5
Risk propensity	GRiPS (Zhang <i>et al.</i> , 2019)	8	0.87	0.88	1–7
Creative self-efficacy	Tierney & Farmer (2002) adapted	6	0.92	0.93	1–5
Prompt relevance	Single item (session-level)	1	-	-	1–5

Notes: Opportunity recognition and creative self-efficacy items were rated on 5-point Likert scales and linearly rescaled to 0–100 using $(x - 1) / 4 \times 100$. Risk propensity was rated on a 7-point Likert scale and rescaled using $(x - 1) / 6 \times 100$. Prompt relevance was analyzed using its original 1-to-5 scale.

4.3.1. Data preparation

To facilitate interpretability, Likert-scale outcomes were linearly rescaled to a 0–100 metric for OR and CSE. For 5-point scales, scores were rescaled using $(x - 1) / 4 \times 100$. For RP, measured on a 7-point scale, scores were rescaled using $(x - 1) / 6 \times 100$. Prompt relevance was analyzed using its original 1-to-5 scale.

4.3.2. Entrepreneurial mindset

Opportunity recognition was measured using the five-item OR Scale developed by Kuckertz *et al.* (2017), which assesses the extent to which students report being alert to and able to identify new business opportunities. Items were rated on a 5-point Likert scale and rescaled to a 0–100 score; α pre = 0.89 and α post = 0.91. RP was measured using the eight-item General Risk Propensity Scale, or GRiPS, developed by Zhang *et al.* (2019). Items were rated on a 7-point Likert scale and rescaled to 0–100; α pre = 0.87 and α post = 0.88.

4.3.3. Creative confidence

Creative self-efficacy was assessed using a six-item scale based on Tierney and Farmer’s (2002) Creative Self-Efficacy measure, with wording adapted for entrepreneurship tasks. Items were rated on a 5-point Likert scale and rescaled to 0–100, α pre = 0.92 and α post = 0.93.

4.3.4. System performance

Prompt relevance was measured after each session using

a single 5-point item that asked students how helpful and well-targeted the AI prompts were for the session's task. Convergence speed was defined as the number of interaction turns required until the normalized composite $Reward_t$ changed by less than 0.08 across three consecutive turns. $Reward_t$ was normalized to the 0 to 1 range before this calculation. Convergence speed was calculated within each session and then averaged across the eight sessions for each student. For fairness across conditions, the same post-hoc rubric was applied to interaction logs from all four conditions, including baseline conditions where reward shaping was not active during the intervention. The full survey instruments are provided in [Appendix A](#), the demographic questionnaire in [Appendix B](#), the interview protocol in [Appendix C](#), illustrative prompt-revision and scaffolding examples in [Appendix D](#), and the reward-shaping rubric in [Appendix E](#).

4.4. Procedure

In Week 0, all participants completed the pre-test battery, including OR, RP, and CSE. The intervention consisted of eight weekly 45-minute AI-supported workshops integrated into regular coursework. Weekly tasks required students to iteratively develop and refine venture ideas in response to realistic market scenarios, such as identifying unmet needs, specifying value propositions, testing assumptions, and proposing low-cost experiments. Sessions followed a standardized task template. The conditions differed only in prompt engineering and adaptation mechanisms. Each workshop followed the same structure:

- (i) Goal setting: Five minutes. Students reviewed their ongoing venture idea and set a session goal.
- (ii) AI-supported tutoring activity: 30 minutes. Students completed the weekly entrepreneurship task using the assigned condition. In the RPCA condition, the full framework was used, including prompt alignment, reward shaping, and confidence tracking. In the static templates condition, students selected from fixed prompts. In the adaptive prompting condition, GPT-4 adjusted task difficulty using the rule-based heuristic described earlier. In the human-AI hybrid condition, students alternated between instructor-written prompts and unguided access to GPT-4.
- (iii) Reflection: 10 minutes. Students completed a prompt-relevance rating and brief, momentary confidence items. System logs captured response latency, interaction turns, follow-up turns, and revision attempts.

In Week 8, participants completed the identical post-test battery. Analyses were conducted using an intention-to-treat approach ($N = 120$). There were no cases of complete

post-test battery missingness. Any sporadic item-level omissions were handled using expectation-maximization imputation. Complete-case analyses yielded substantively identical conclusions. Qualitative data from post-study interviews ($n = 8$) were collected but are not reported in the present manuscript.

4.5. Ablation study

To examine the contribution of RPCA's components, a within-subject ablation study was conducted during Weeks 7 and 8 within the RPCA condition. Each RPCA participant completed four matched-difficulty ideation tasks across four configurations: full RPCA, RPCA without the reward-shaping engine, RPCA without the confidence tracking layer, and a baseline GPT-4 with a static prompt. The four configurations were counterbalanced using a Latin square design to reduce order effects. Task difficulty was matched through pilot testing.

After each ablation task, students completed brief post-task ratings of OR, CSE, and prompt relevance. These outcomes were used to test whether removing reward shaping or confidence tracking reduced performance relative to the full RPCA configuration.

4.6. Data analysis

For H1 and H2, mixed-design analyses of variance (ANOVAs) were conducted with time as the within-subjects factor and condition as the between-subjects factor to test pre-to-post changes in OR and RP across the four conditions. For H3, analysis of covariance (ANCOVA) was used to compare post-test CSE across conditions while controlling for pre-test CSE. For H4, one-way ANOVAs were used to compare mean prompt relevance and convergence speed across conditions, followed by Tukey-adjusted post-hoc comparisons.

For H5 and H6, repeated-measures ANOVAs were conducted on the ablation data, with configuration as the within-subject factor. Planned contrasts compared the full RPCA configuration with the no-reward-shaping and no-confidence-tracking configurations. Holm adjustment was applied to control for multiple comparisons.

Assumption checks were conducted before hypothesis testing. Normality and homogeneity of variance were examined for between-group analyses. For repeated-measures analyses, sphericity was assessed, and Greenhouse-Geisser corrections were applied when necessary. Statistical significance was evaluated at $p < 0.05$.

4.7. Ethics, consent, and data privacy

This study was approved by the Institutional Review Board of Raffles International College, Bangkok. All participants

provided informed consent before participation. Students were informed that participation was voluntary, that they could withdraw at any time, and that participation or withdrawal would not affect their course grades.

All interaction logs were anonymized using StudentID codes only. No directly identifiable personal information was entered into the GPT-4 API. The data were used only for research purposes and stored on an encrypted server accessible only to the research team.

5. Results

5.1. Pre- and post-test scores

Pre-test scores did not differ significantly between conditions on any measure (all $p > 0.41$), indicating comparable baseline scores across the four randomized conditions (Table 2). Assumption checks did not indicate serious violations of normality or homogeneity of variance. For repeated-measures analyses in which sphericity was violated, Greenhouse–Geisser corrections were applied.

Students in the RPCA condition showed significantly larger gains in both OR and RP than those in all baseline

conditions, providing evidence for H1 and H2.

For OR (H1), RPCA students improved from a mean of 61.2 (SD = 12.4) on the pre-test to 82.8 (SD = 10.1) on the post-test (+21.6 points), whereas the three baseline conditions gained only 10.7–11.4 points (post-test range: 71.5–73.2). A mixed-design ANOVA revealed a significant condition $\times$ time interaction ($F[3, 116] = 18.42, p < 0.001, \eta^2 = 0.323$) (Table 2). Tukey post-hoc tests confirmed that RPCA significantly outperformed all baselines (all $p < 0.001$).

For RP (H2), RPCA students gained +14.4 points (pre-mean = 54.3, SD = 11.8; post-mean = 68.7, SD = 10.5), compared with +6.3 to +7.7 points in the baseline conditions. The condition $\times$ time interaction was significant ($F[3, 116] = 11.68, p < 0.001, \eta^2 = 0.232$), with RPCA significantly higher than all baselines in post-hoc tests (all $p < 0.001$). This suggests that RPCA's explicit emphasis on calculated risk-taking supports shifts in students' comfort with uncertainty, aligning with recent evidence that entrepreneurship education can meaningfully shape risk-related attitudes (Hou *et al.*, 2022; López-Muñoz *et al.*, 2023).

Table 2. Pre- and post-test means and gain scores by condition

Condition	OR pre	OR post	OR gain	RP pre	RP post	RP gain	CSE pre	CSE post	CSE gain
RPCA	61.2 (12.4)	82.8 (10.1)	+21.6	54.3 (11.8)	68.7 (10.5)	+14.4	59.5 (13.0)	77.5 (11.5)	+18.0
Static templates	60.8 (11.9)	71.5 (12.3)	+10.7	53.9 (12.1)	60.2 (11.4)	+6.3	59.0 (12.8)	51.0 (12.5)	−8.0
Adaptive prompting	62.1 (12.6)	73.2 (11.8)	+11.1	55.1 (11.5)	61.8 (12.0)	+6.7	58.5 (13.5)	59.8 (13.0)	+1.3
Human–AI hybrid	61.5 (11.7)	72.9 (12.0)	+11.4	54.7 (12.3)	62.4 (11.7)	+7.7	58.8 (12.9)	64.8 (12.6)	+6.0

Notes: Values are means (SD) on rescaled 0–100 scores. Gains are post–pre differences. OR = opportunity recognition; RP = risk propensity; CSE = creative self-efficacy. CSE adjusted post-test means (analysis of covariance controlling for pre-test CSE) are reported in Table 3. Groups were equal in size ($n = 30$ per condition).

Abbreviations: AI: Artificial intelligence; CSE: Creative self-efficacy; OR: Opportunity recognition; RP: Risk propensity; RPCA: Reciprocal Prompt Co-Adaptation; SD: Standard deviation.

Table 3. Adjusted mean creative self-efficacy post-test scores by condition

Condition	n	Adjusted mean	Standard error	95% confidence interval	Pairwise vs. RPCA (Bonferroni-adjusted p)
RPCA	30	78.0	2.0	[74.0, 82.0]	–
Static templates	30	50.3	2.2	[46.0, 54.6]	<0.001
Adaptive prompting	30	59.2	2.0	[55.2, 63.1]	<0.001
Human–AI hybrid	30	65.2	2.2	[60.9, 69.5]	<0.001

Notes: Adjusted means and standard errors are from analysis of covariance on the rescaled 0–100 CSE score. Main effect of condition: $F(3, 115) = 38.42, p < 0.001, \eta^2 = 0.50$. Pre-test covariate: $F(1, 115) = 12.67, p < 0.001, \eta^2 = 0.10$.

Abbreviations: AI: Artificial intelligence; RPCA: Reciprocal Prompt Co-Adaptation.

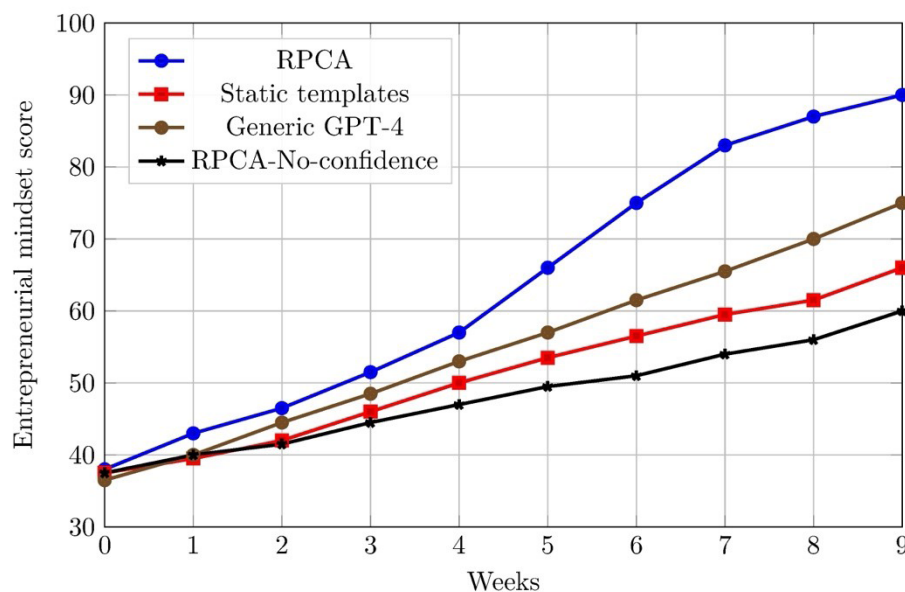


Figure 2. Entrepreneurial mindset development trajectories across experimental conditions over eight weeks
Abbreviation: RPCA: Reciprocal Prompt Co-Adaptation.

Figure 2 shows the entrepreneurial mindset trajectories over the eight weeks, addressing RQ1 at the process level. The RPCA curve exhibits an S-shaped pattern with a pronounced mid-semester increase, while baselines display more linear and gradual improvements.

Reciprocal Prompt Co-Adaptation also yielded stronger gains in CSE than the baseline approaches, addressing H3. On the rescaled 0 to 100 CSE score, RPCA students improved from a mean of 59.5 (SD = 13.0) at pre-test to 77.5 (SD = 11.5) at post-test, a gain of 18.0 points. The baseline conditions showed smaller changes: static templates: +8.0; adaptive prompting: +1.3; and human-AI hybrid: +6.0.

- An ANCOVA controlling for pre-test CSE showed a significant effect of condition on post-test CSE ($F[3, 115] = 38.42, p < 0.001, \eta^2 = 0.50$) (Table 3). The pre-test covariate was also significant ($F[1, 115] = 12.67, p < 0.001, \eta^2 = 0.10$). Bonferroni-adjusted comparisons of adjusted means indicated that RPCA scored higher than static templates, adaptive prompting, and human-AI hybrid (all $p < 0.001$).
- These results are consistent with the expectation that structured reflection and iterative feedback can support students' CSE. The process indicators also suggest that RPCA offered a more targeted interaction environment, as reflected in higher session-level prompt relevance and faster convergence to stable prompts. However, these process indicators should be

interpreted as evidence of interaction quality rather than as direct proof of AI-side learning.

- Taken together, these findings provide convergent support for RQ1 and H1–H3, showing that RPCA was associated with larger gains in OR, RP, and CSE than the comparison conditions.

Prompt relevance ratings were recorded after each session on a 1–5 scale and averaged across eight sessions for each participant. A one-way ANOVA showed a significant main effect of condition on mean prompt relevance ($F[3, 116] = 48.67, p < 0.001, \eta^2 = 0.56$) (Table 4). Tukey-adjusted post-hoc tests indicated that RPCA (mean = 4.52, SD = 0.48) outperformed static templates (mean = 2.68, SD = 0.82), adaptive prompting (mean = 3.45, SD = 0.71), and human-AI hybrid (mean = 3.91, SD = 0.66) (all $p < 0.001$).

Convergence speed was operationalized as the number of interaction turns required for the normalized composite Reward_t to change by less than 0.08 across three consecutive turns. Although reward shaping was active only in the RPCA condition during the intervention, Reward_t was computed post-hoc from interaction logs using the same rubric across all conditions to support fair process comparison. For each participant, convergence speed was computed within each session and then averaged across the eight sessions. A one-way ANOVA indicated a significant effect of condition on convergence speed ($F[3, 116] = 52.11, p < 0.001, \eta^2 = 0.57$) (Table 5).

RPCA required fewer turns to converge (mean = 3.40, SD = 1.20) than the three baseline conditions. Together, these findings support H4.

Table 4. Mean prompt relevance ratings by condition

Condition	<i>n</i>	Mean	SD	95% CI
RPCA	30	4.52	0.48	[4.34, 4.70]
Static templates	30	2.68	0.82	[2.37, 2.99]
Adaptive prompting	30	3.45	0.71	[3.18, 3.72]
Human–AI hybrid	30	3.91	0.66	[3.66, 4.16]

Notes: Ratings on a 5-point scale. ANOVA: $F(3, 116) = 48.67, p < 0.001, \eta^2 = 0.56$. Tukey-adjusted post-hoc comparisons were conducted following the omnibus ANOVA.

Abbreviations: AI: Artificial intelligence; ANOVA: Analysis of variance; CI: Confidence interval; RPCA: Reciprocal Prompt Co-Adaptation; SD: Standard deviation.

Table 5. Mean convergence speed (turns to stability) by condition

Condition	<i>n</i>	Mean	Standard deviation
RPCA	30	3.40	1.20
Static templates	30	7.80	2.10
Adaptive prompting	30	6.20	1.80
Human–AI hybrid	30	5.50	1.60

Notes: Lower values indicate faster convergence. Convergence speed was calculated as the mean number of turns required for normalized Reward_t to stabilize within the 0.08 threshold across three consecutive turns.

ANOVA: $F(3, 116) = 52.11, p < 0.001, \eta^2 = 0.57$.

Abbreviations: AI: Artificial intelligence; ANOVA: Analysis of variance; RPCA: Reciprocal Prompt Co-Adaptation.

Figure 3 illustrates the relationship between CreativityScore and RiskScore in RPCA student responses. The positive association suggests that, within the rubric-guided scoring procedure, responses with higher creativity scores also tended to receive higher calculated-risk scores. This pattern is consistent with the intended balance between creative exploration and justified risk-taking. However, these model-assisted scores should be interpreted as process indicators rather than external assessments of idea quality.

These process-level results support H4, indicating that RPCA improved prompt relevance and convergence efficiency relative to baseline prompting approaches.

5.2. Exploratory quality checks for model-assisted indicators

To address the transparency and stability of the model-assisted process indicators, additional quality checks were conducted on the textual confidence estimates, reward scores, and confidence estimates used in the RPCA condition. These analyses were exploratory and intended to assess whether the rubric-guided indicators aligned with researcher judgments and student-reported confidence.

First, the research team manually reviewed approximately 20% of the logged student responses used for textual confidence coding ($n = 1,247$). Agreement between GPT-4-based TextSignal_t scores and researcher judgments was good (intraclass correlation coefficient [ICC] = 0.81; 95% CI: [0.79, 0.83]; $p < 0.001$). This suggests that the fixed rubric for extracting textual confidence indicators was applied consistently.

Second, approximately 15% of student responses ($n = 936$) were manually reviewed for reward-scoring quality. Agreement between GPT-4-based scores and researcher judgments was acceptable to good for CreativityScore (ICC = 0.79, 95% CI: [0.76, 0.81]), RiskScore (ICC = 0.77, 95% CI: [0.74, 0.80]), and the composite Reward_t (ICC = 0.82, 95% CI: [0.80, 0.84]). These results support the use of rubric-guided reward scores as process indicators, although they should not be interpreted as external ground-truth measures of entrepreneurial performance.

Third, the confidence estimate C_t showed a moderate positive correlation with students' momentary confidence reports ($r = 0.58, p < 0.001$). Mean C_t was also positively associated with CSE gains ($r = 0.47, p < 0.001$). These correlations provide preliminary convergent evidence that the confidence tracking layer captured variation in students' perceived creative confidence.

Finally, session-level C_t showed moderate stability across the eight intervention sessions, with an ICC = 0.69 for single measures and a Cronbach's $\alpha = 0.84$ across the eight session-level estimates. This pattern suggests that the confidence estimates were sufficiently stable for adaptive scaffolding purposes while still allowing session-to-session variation in learner state.

5.3. Ablation study

To assess the contribution of RPCA's individual components and address RQ3, a within-subject ablation study was embedded in the RPCA cohort during the final two weeks, Weeks 7 and 8. Each RPCA participant completed four brief ideation tasks, matched in format and difficulty, across four system-level configurations. The order of the four configurations was counterbalanced using a Latin

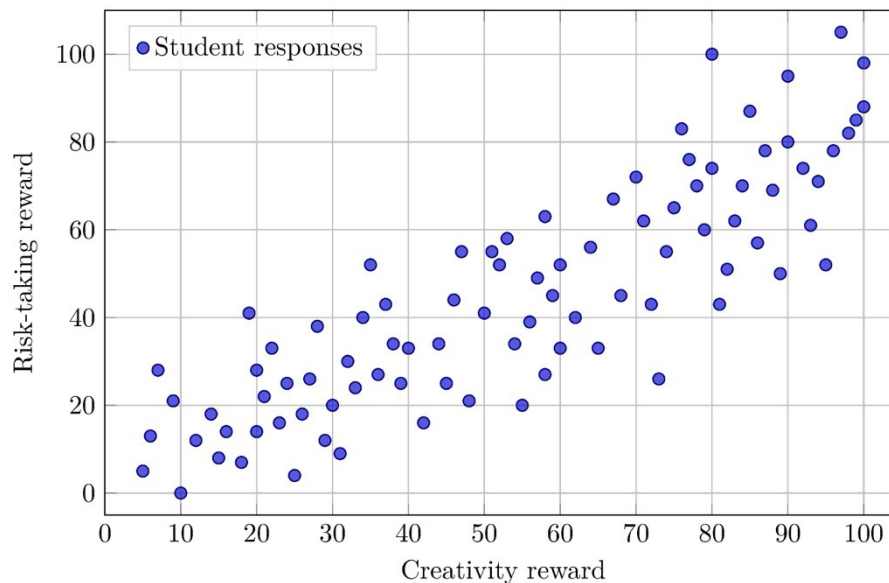


Figure 3. Relationship between creativity and risk-taking rewards in student responses

square design, and task difficulty was matched through pilot testing.

The four configurations were as follows: (i) Full RPCA: all components active; (ii) no reward shaping: prompt refinement and confidence tracking retained, but the entrepreneurship-specific composite Reward_i was not used to guide feedback or prompt updates; (iii) no confidence tracking: reward shaping and prompt refinement retained, but coaching did not use the confidence estimate C_i ; and (iv) baseline GPT-4: a single static system prompt with no RPCA mechanisms.

After each ideation task, participants completed brief post-task self-ratings of OR and CSE and rated prompt relevance on a 1–5 scale. For each configuration, post-task OR and CSE scores were analyzed on a 0–100 scale, and prompt relevance was analyzed on its original 1–5 scale. Table 6 reports descriptive means and standard deviations for the within-subject comparisons ($n = 30$).

Repeated-measures ANOVAs indicated significant effects of configuration on OR, CSE, and prompt relevance. Holm-adjusted planned comparisons between full RPCA and the removal of reward shaping showed that removing reward shaping reduced OR ($t[29] = 5.82, p < 0.001$) and CSE ($t[29] = 4.31, p < 0.001$).

For H5 (reward shaping), relative to full RPCA (OR = 84.2 and CSE = 74.6), the no reward shaping configuration showed lower OR (mean = 72.1) and CSE (mean = 69.3). Holm-adjusted planned comparisons confirmed these reductions, OR: $t(29) = 5.82, p < 0.001$; CSE: $t(29) = 4.31,$

$p < 0.001$. This suggests that the entrepreneurship-specific reward signal helped keep the interaction aligned with mindset-development goals.

Table 6. Ablation study: Post-task performance

Configuration	OR score (0–100)	CSE score (0–100)	Prompt relevance (1–5)
Full RPCA	84.2 (9.8)	74.6 (10.3)	4.38 (0.49)
No reward shaping	72.1 (11.4)	69.3 (11.7)	4.01 (0.62)
No confidence tracking	76.8 (10.7)	65.9 (12.1)	3.87 (0.71)
Baseline GPT-4 (static)	68.4 (12.3)	63.2 (13.0)	3.45 (0.83)

Notes: Within-subject comparisons inside the RPCA cohort ($n = 30$). OR/CSE scores are 0–100; prompt relevance is 1–5. Higher values indicate better performance. Omnibus effects were tested using repeated-measures ANOVA with Greenhouse–Geisser correction (OR: $F[3, 87] = 28.41, p < 0.001, \eta^2 = 0.49$; CSE: $F[3, 87] = 22.67, p < 0.001, \eta^2 = 0.44$; prompt relevance: $F[3, 87] = 19.83, p < 0.001, \eta^2 = 0.41$). Holm-adjusted planned contrasts versus full RPCA are reported in the text.

Abbreviations: CSE: Creative self-efficacy; OR: Opportunity recognition; RPCA: Reciprocal Prompt Co-Adaptation.

For H6 (confidence tracking), disabling the confidence tracking layer primarily reduced CSE (mean = 65.9) and perceived prompt quality (prompt relevance mean = 3.87), compared with full RPCA (CSE = 74.6 and prompt relevance = 4.38). Holm-adjusted planned comparisons confirmed these reductions, CSE: $t(29) = 5.67, p < 0.001$;

prompt relevance: $t(29) = 4.18, p < 0.001$. This pattern is consistent with confidence-aware coaching serving as a scaffolding mechanism to sustain perceived task manageability and prompt quality during open-ended ideation.

To situate these configurations relative to the main experimental conditions, Table 7 summarizes the key design differences.

These patterns provide preliminary support for H5 and H6 and suggest that reward shaping and confidence tracking contributed differently to the observed RPCA effects. Reward shaping appeared to support alignment with entrepreneurial learning goals, whereas confidence-aware coaching appeared to support students' perceived manageability and prompt relevance during complex ideation tasks.

5.4. Summary of hypothesis testing

The experiment provided convergent support for all six hypotheses. The full RPCA framework produced significantly larger gains in OR (H1), RP (H2), and CSE (H3) than the three baseline conditions. It also achieved higher prompt relevance ratings and required fewer interaction turns to reach convergence (H4; Tables 4 and 5). Ablation analyses further supported the contribution of the two core adaptive components. Removing the reward-shaping engine reduced OR and CSE (H5),

whereas removing the confidence tracking layer reduced CSE and perceived prompt quality (H6). Table 8 presents a consolidated overview.

Taken together, these results provide initial evidence that feedback-driven prompt co-adaptation, supported by domain-specific reward shaping and confidence-aware scaffolding, can improve student-side entrepreneurial mindset outcomes and the quality of interaction processes when implemented with an off-the-shelf LLM. The results directly address RQ1 to RQ3 and support H1 to H6, while keeping the empirical interpretation centered on student learning and interaction-level adaptation rather than AI-side learning.

6. Discussion

Building on the empirical support for RQ1–RQ3 and H1–H6, this section discusses the contribution, limitations, broader applications, and ethical implications of the current RPCA implementation. The quantitative results indicate that RPCA outperformed the baseline conditions on student-side entrepreneurial mindset outcomes, including OR, RP, and CSE, as well as on interaction process indicators such as prompt relevance and convergence speed. The within-cohort ablation analyses further suggest that reward shaping and confidence tracking contributed differently to the observed effects within the RPCA condition.

Table 7. Comparison of experimental conditions

Condition / Configuration	Prompt refinement	Reward signal	Confidence adaptation	Key difference from full RPCA
RPCA (full)	Student + AI	Entrepreneurship-specific rubric	Real-time tracking	–
Static templates	None, fixed prompts	None	None	No refinement or adaptation
Adaptive prompting	AI only (performance-based)	Generic correctness and engagement	None	One-sided adaptation, no domain reward or confidence tracking
Human–AI hybrid	Student only (manual)	None	None	The student has full control, no automated RPCA mechanisms
No reward shaping	Student + AI	None, no composite Reward _i	Real-time tracking	Removes domain-specific reward signal
No confidence tracking	Student + AI	Entrepreneurship-specific rubric	None	Removes confidence-aware coaching
Baseline GPT-4 (static)	None (single prompt)	None	None	No RPCA mechanisms at all

Notes: Static templates (main experiment) refers to choosing from a fixed prompt menu (templates do not change across sessions); Baseline GPT-4 (static) (ablation) refers to a single fixed system prompt with no RPCA mechanisms. Abbreviations: AI: Artificial intelligence; RPCA: Reciprocal Prompt Co-Adaptation.

Table 8. Summary of hypothesis tests

Hypothesis	Statement	Result	Primary evidence
H1	RPCA > baselines on opportunity recognition (OR) gains	Supported	Table 2; $F(3, 116) = 18.42, p < 0.001, \eta^2 = 0.323$
H2	RPCA > baselines on risk propensity (RP) gains	Supported	Table 2; $F(3, 116) = 11.68, p < 0.001, \eta^2 = 0.232$
H3	RPCA > baselines on creative self-efficacy (CSE) gains	Supported	Table 3; ANCOVA (pre-test CSE covariate), $F(3, 115) = 38.42, p < 0.001, \eta^2 = .50$
H4	RPCA yields higher prompt relevance and faster convergence	Supported	Table 4; $F(3, 116) = 48.67, p < 0.001, \eta^2 = 0.56$; Table 5; $F(3, 116) = 52.11, p < 0.001, \eta^2 = 0.57$
H5	Removing the reward-shaping engine reduces OR and CSE task performance	Supported	Table 6; RM-ANOVA: OR: $F(3, 87) = 28.41, p < 0.001, \eta^2 = 0.49$; CSE: $F(3, 87) = 22.67, p < 0.001, \eta^2 = 0.44$; Holm contrasts vs. Full RPCA: no reward shaping: OR: $t(29) = 5.82, p < 0.001$; CSE: $t(29) = 4.31, p < 0.001$
H6	Removing confidence tracking reduces CSE and perceived prompt quality	Supported	Table 6; RM-ANOVA: CSE: $F(3, 87) = 22.67, p < 0.001, \eta^2 = 0.44$; prompt relevance: $F(3, 87) = 19.83, p < 0.001, \eta^2 = 0.41$; Holm contrasts vs. Full RPCA: no confidence tracking: CSE: $t(29) = 5.67, p < 0.001$; prompt relevance: $t(29) = 4.18, p < 0.001$

These findings should be interpreted as evidence of student learning, supported by interaction-level prompt adaptation, rather than as evidence that GPT-4 learned in the human sense or acquired new knowledge during the intervention. In RPCA, adaptation occurs through prompt revision, rubric-guided feedback, and confidence-aware scaffolding. The underlying model parameters remain unchanged. Therefore, the contribution of the framework lies in designing a structured, feedback-driven tutoring protocol that allows students to influence prompt and scaffolding decisions while maintaining an empirical focus on student outcomes.

6.1. Interpretation of findings

The findings suggest that feedback-driven prompt co-adaptation may be useful for supporting entrepreneurial mindset development in open-ended learning contexts. Compared with static templates and one-sided adaptive prompting, RPCA enabled students to assess prompt relevance, offer brief reflections, and inform subsequent prompt revisions. This interpretation aligns with recent work that frames prompt engineering as an emerging educational skill, in which learners need to understand how task framing, constraints, and prompt design shape AI-supported reasoning and output quality (Federiakin *et al.*, 2024; Qian, 2025). In the present study, this may have helped students engage more actively with the ideation process rather than simply responding to pre-written or AI-generated prompts.

More broadly, recent meta-analytic evidence suggests that ChatGPT-based interventions can support learning

outcomes when embedded in educational contexts, although effects vary by outcome type and implementation design (Deng *et al.*, 2025).

The stronger gains in OR and RP suggest that the entrepreneurship-specific reward structure may have helped keep the interaction aligned with the learning goals of the course. Instead of rewarding only fluency or task completion, the system emphasized creativity and calculated risk-taking. This design appears to have encouraged students to generate more varied ideas, consider uncertainty, and justify risk-related decisions. However, this interpretation should be approached with caution because Reward_t was a model-assisted process indicator rather than an external assessment of entrepreneurial competence.

The gains in CSE may be partly explained by the confidence-aware scaffolding mechanism. This interpretation is also consistent with prior creativity research showing that structured self-assessment and reflection can support students' creative self-efficacy and creative thinking processes (Z. Yan *et al.*, 2022). Students with lower estimated confidence received more concrete examples and step-by-step guidance, while students with higher estimated confidence received more open-ended challenges. This adaptive pattern may have helped maintain a balance between support and productive difficulty. The exploratory checks reported in Section 5 provide preliminary evidence that C_t was related to students' momentary confidence and CSE gains. At the same time, these checks do not make C_t a fully validated psychological measure. Rather, they suggest that the confidence estimate

behaved in a reasonable and interpretable way within this classroom prototype.

The ablation findings further clarify the role of each component. Removing reward shaping reduced OR and CSE, suggesting that domain-specific feedback criteria helped anchor the interaction to entrepreneurial learning goals. Removing confidence tracking reduced CSE and prompt relevance, suggesting that confidence-aware scaffolding helped students experience the prompts as more manageable and better targeted. These findings provide preliminary support for the component logic of RPCA, while still requiring replication with larger samples and longer interventions.

6.2. Limitations of the Reciprocal Prompt Co-Adaptation system

Several limitations should be considered when interpreting the findings. At the contextual level, the study was conducted with Thai undergraduate business students over an eight-week intervention. The robustness and durability of the effects in other disciplines, age groups, institutional contexts, and cultural settings remain uncertain. Future research should examine longer interventions, delayed post-tests, and cross-institutional replications to determine whether the observed gains are sustained over time.

The embedded ablation study also has limitations. Although the order of configurations was counterbalanced using a Latin square design and the tasks were matched through pilot testing, the ablation was conducted during the final two weeks of the intervention. Students had already gained experience with RPCA by that point, so residual learning, familiarity, or carryover effects across configurations cannot be fully ruled out. Future studies could use separate participant groups, longer washout periods, or fully independent ablation experiments to strengthen causal claims about individual components.

A further contextual constraint concerns interaction language. All student-GPT-4 interactions were conducted in English rather than Thai. For Thai students, operating in a second language may have increased cognitive load, limited nuance in articulating ideas, and influenced their willingness to take linguistic and conceptual risks. Language choice may also affect how the model interprets student intent, politeness strategies, expressions of uncertainty, and confidence signals. Future work should examine bilingual or Thai-first deployments and test whether language moderates the effects of RPCA on OR, RP, and CSE.

At the system level, the framework depends on a single proprietary LLM, GPT-4. This reliance may inherit biases from the model's training data and alignment process,

including dominant business narratives, culturally preferred communication styles, or conventional assumptions about entrepreneurship (Alvero *et al.*, 2024). Such biases could constrain the range of entrepreneurial ideas students explore or steer them toward familiar opportunity framings. Future implementations should compare different models, include localized examples, and examine whether model choice affects the diversity and cultural relevance of student ideas.

The reward-shaping component also requires cautious interpretation. RPCA uses rubric-guided GPT-4 scoring to estimate creativity and calculated risk-taking. Although exploratory quality checks showed acceptable agreement between GPT-4-based scores and researcher judgments, these scores remain constrained by the rubric design and the model's interpretive stability. Reward_t should therefore be treated as a process indicator, not as an external ground-truth measure of entrepreneurial ability or idea quality. Future studies should combine model-based scoring with independent expert ratings, peer assessments, behavioral indicators, or performance-based outcomes such as prototype quality, customer discovery plans, or pitch evaluations. This caution is consistent with broader concerns in AI ethics in education, where model-assisted indicators should be treated as decision-support tools rather than authoritative judgments of student ability (Akgun & Greenhow, 2022; Y. Yan *et al.*, 2025).

The confidence tracking layer has similar limitations. C_t integrates textual, behavioral, and self-report signals through an a priori rule-based weighting scheme rather than a trained multimodal estimator. This design improves transparency and classroom deployability, but it may be sensitive to contextual noise. Network latency, device differences, writing fluency, cultural communication norms, and students' willingness to express uncertainty may all influence confidence estimation. The exploratory checks in this study provide preliminary support for consistency and convergent association, but further validation is needed before C_t can be treated as a robust learner-state measure. This is especially important because learning analytics indicators can increase learner visibility but may also be misinterpreted if their uncertainty, context dependence, and intended use are not made explicit (Gedrimiene *et al.*, 2023; Susnjak *et al.*, 2022).

Finally, deploying GPT-4-based tutoring in real courses raises organizational and regulatory constraints. Institutions must consider Application Programming Interface costs, technical infrastructure, data protection, consent procedures, and the governance of learning analytics. Fine-grained interaction logs can contain students' ideas, reflections, expressions of uncertainty,

and confidence reports. Long-term deployment of RPCA, therefore, requires clear procedures for data minimization, anonymization, storage, access control, and auditing, especially when similar systems are integrated into institutional learning management systems. These concerns are consistent with prior reviews of privacy and data protection in learning analytics and with emerging regulatory discussions around generative AI governance, cybersecurity, and institutional responsibility (Liu & Khalil, 2023; Novelli *et al.*, 2024).

6.3. Potential application scenarios beyond entrepreneurial tutoring

Although the empirical evaluation focused on entrepreneurship education, the RPCA design may be adaptable to other domains that require open-ended reasoning, iterative feedback, and learner agency. Its core principles are prompt refinement, domain-specific feedback criteria, and confidence-aware scaffolding. These principles are not limited to entrepreneurship, but the reward criteria and scaffolding rules would need to be redesigned for each learning context.

In academic writing and communication courses, prompts could be refined to focus on argument quality, rhetorical flexibility, audience awareness, and use of evidence. Reward criteria could emphasize clarity, coherence, originality, and responsiveness to feedback. In such settings, RPCA could help students reflect on how different prompt framings shape the development of their arguments. This would position the system as a structured writing scaffold rather than a replacement for student authorship or teacher feedback.

In professional training contexts, RPCA could support scenario-based learning in fields such as management, negotiation, public speaking, or leadership development. This extension is relevant to business and management education, where recent work has identified both opportunities and tensions in the use of ChatGPT for classroom learning and professional skill development (Valcea *et al.*, 2024). Domain-specific reward criteria could be designed around situational judgment, stakeholder sensitivity, ethical reasoning, or strategic clarity. Confidence-aware scaffolding could then adjust the level of challenge based on whether learners appear ready for more complex scenarios or require more guided practice.

Cross-cultural learning environments represent another possible application area. Because the prompt-alignment module can incorporate localized examples and student feedback, RPCA may help adapt AI-mediated instruction to diverse student cohorts. However, this potential also requires caution. LLMs may reproduce dominant cultural

assumptions unless explicitly counterbalanced. Recent studies similarly show that LLM outputs can reflect cultural dominance, social values, and culturally patterned bias, making cultural alignment and localized review important for cross-cultural educational deployment (Dokic *et al.*, 2025; Tao *et al.*, 2024). Future deployments should therefore combine prompt co-adaptation with localized content development, cultural bias checks, and stakeholder review.

Across these possible applications, the framework should not be interpreted as implying that the AI is a human-like learning partner. Its value lies in making prompt evolution visible, inspectable, and responsive to student input. This may strengthen learner agency and metacognitive awareness while keeping human judgment central to the learning process.

6.4. Ethical considerations in the Reciprocal Prompt Co-Adaptation system

Reciprocal Prompt Co-Adaptation's capacity to shape learners' confidence, risk orientation, and feedback engagement raises ethical questions related to autonomy, transparency, and pedagogical influence. Systems that adapt prompts and challenges based on learner-state estimates must distinguish legitimate educational scaffolding from manipulative steering. This concern echoes broader discussions of AI ethics in education, especially around transparency, autonomy, fairness, and the appropriate boundaries of automated pedagogical influence (Akgun & Greenhow, 2022; Y. Yan *et al.*, 2025). This is especially important in domains such as entrepreneurship, where the system may influence students' attitudes toward uncertainty, risk, and opportunity.

In the present study, RPCA was limited to entrepreneurship education and used for formative learning support rather than high-stakes assessment. Even so, similar techniques could be adapted to more sensitive domains, including political attitudes, moral reasoning, or personal decision-making. Responsible use, therefore, requires clear boundaries around appropriate learning contexts, transparent communication of educational objectives, and safeguards against persuasive uses that undermine learner agency.

Data privacy is another central concern. RPCA relies on interaction logs, student reflections, confidence reports, and model-assisted indicators. Students may not always understand how such data are interpreted or used. Learning analytics research similarly emphasizes that students' awareness, expectations, and consent should be considered when data traces are used to guide educational decisions (Gedrimiene *et al.*, 2023; Liu & Khalil, 2023). In future

deployments, students should be informed about what data is collected, how confidence and reward indicators are generated, who can access them, and whether they affect assessment. These indicators should not be used for summative grading unless they have undergone stronger validation and are accompanied by appeal or review procedures.

Future iterations of RPCA should incorporate stronger technical and governance safeguards. On the technical side, aggregation, pseudonymization, and privacy-preserving analytics could reduce the risk of exposing sensitive learner traces. On the governance side, explainability features could allow students and instructors to inspect how prompts evolved, why a specific challenge was selected, or how a confidence estimate was derived. Continuous monitoring is also needed to detect unintended reinforcement of narrow business ideologies, cultural stereotypes, or discriminatory patterns.

7. Conclusion

This paper examined RPCA, a structured human–AI scaffolding framework designed to support entrepreneurial mindset development and CSE through feedback-driven prompt refinement with GPT-4. Rather than treating prompts as static scripts or relying only on one-way AI adaptation, RPCA allows students to evaluate prompt relevance, provide reflections, and influence subsequent prompt revisions. The system then uses domain-specific reward criteria and confidence-aware scaffolding rules to adjust the tutoring interaction. Importantly, RPCA does not assume that GPT-4 learns in the human sense or updates its model parameters. The adaptive element lies in the interaction protocol, not in model-side learning.

The eight-week classroom study across three Thai universities provides initial evidence for the effectiveness of this approach. Relative to static templates, simple adaptive prompting, and human–AI hybrid prompting, RPCA produced larger gains in OR, RP, and CSE. It also achieved higher prompt-relevance ratings and faster convergence to stable prompts. The ablation analyses further suggested that the reward-shaping engine and confidence tracking layer contributed differently to the framework's effects. Removing reward shaping reduced OR and CSE, whereas removing confidence tracking reduced CSE and perceived prompt quality.

Beyond these performance effects, RPCA offers a conceptual contribution by reframing LLM-based tutoring as configurable dialogic scaffolding rather than as AI-side learning or peer-like collaboration. Making prompt evolution explicit and responsive to student feedback may strengthen learners' agency and metacognitive awareness

about how different framings shape idea generation and evaluation. In entrepreneurship education, this may support richer OR processes, more balanced attitudes toward risk, and a greater willingness to iterate on imperfect ideas.

Future research should extend this work in several directions. Larger and more diverse samples, longer intervention periods, and delayed post-tests are needed to examine the robustness and durability of RPCA effects. Further validation is also needed for the reward and confidence indicators, ideally through external expert ratings, behavioral performance measures, and cross-cultural replications. Technically, future versions should integrate stronger safeguards for explainability, privacy, and governance. As generative AI becomes more deeply embedded in education, approaches such as RPCA illustrate how learner-centered, transparent, and feedback-driven interaction designs can support human creativity and judgment without overstating the learning capacity of the AI system.

Acknowledgments

None.

Funding

None.

Conflict of interest

The authors declare that they have no competing interests.

Author contributions

Conceptualization: Qinjie Shen, Panupong Pituk Sung

Data curation: Qinjie Shen

Formal analysis: Qinjie Shen

Investigation: All authors

Methodology: Qinjie Shen, Panupong Pituk Sung

Project administration: Qinjie Shen, Panupong Pituk Sung

Supervision: Panupong Pituk Sung

Writing—original draft: Qinjie Shen

Writing—review & editing: All authors

All authors have read and agreed to the submitted version of the manuscript.

Ethics approval and consent to participate

This study was approved by the Institutional Review Board of Raffles International College Bangkok (Ethics Exemption Number: RICB-ETH-2025-0015). The approval covered the classroom-based study involving undergraduate students, including the collection of pre-test and post-test survey data, student interaction logs, prompt-relevance ratings, momentary confidence ratings,

and anonymized learning-process indicators. Where required, permission to conduct the classroom-based study was obtained from the participating institutions or course instructors. All participants provided informed consent before participation. Students were informed that participation was voluntary, that they could withdraw at any time, and that participation or withdrawal would not affect their course grades.

Consent for publication

Not applicable. The manuscript does not contain identifiable personal data, images, or individual case information requiring consent for publication. All reported data were anonymized and presented in aggregate form.

Availability of data

The datasets generated and analyzed in the current study are not publicly available due to participant privacy and institutional data protection considerations. Anonymized data may be made available from the corresponding authors, Shen Qinjie (SHENQINJIE@rafflesinternationalcollege.ac.th) and Panupong Pituk Sung (panupongp@raffles.education), upon reasonable request and subject to ethical and institutional approval.

Further disclosure

During the preparation of this manuscript and study, the authors used GPT-4 as part of the research intervention and model-assisted process analysis described in the Method section. GPT-4 was used to provide AI-supported tutoring prompts, support rubric-guided scoring of student responses, assist with textual confidence feature extraction, and implement GPT-4 edit-based prompt revision within the RPCA framework. The authors reviewed, checked, and interpreted the outputs and take full responsibility for the originality, validity, and integrity of the manuscript and the reported findings.

References

- Abulela, M. A. (2023). Development and initial validation of a creative self-efficacy scale for undergraduates: Categorical confirmatory factor analysis and multidimensional item response theory. *Frontiers in Education*, 8. <https://doi.org/10.3389/feduc.2023.1306532>
- Akgun, S., & Greenhow, C. (2022). Artificial intelligence in education: Addressing ethical challenges in K-12 settings. *AI and Ethics*, 2(3), 431–440. <https://doi.org/10.1007/s43681-021-00096-7>
- Alvero, A. J., Lee, J., Regla-Vargas, A., Kizilcec, R. F., Joachims, T., & Antonio, A. L. (2024). Large language models, social demography, and hegemony: Comparing authorship in human and synthetic text. *Journal of Big Data*, 11(1), 138. <https://doi.org/10.1186/s40537-024-00986-7>
- Chaudhari, S., Aggarwal, P., Murahari, V., Rajpurohit, T., Kalyan, A., Narasimhan, K., Deshpande, A., & Castro da Silva, B. (2025). RLHF deciphered: A critical analysis of reinforcement learning from human feedback for LLMs. *ACM Computing Surveys*, 58(2), 1–37. <https://doi.org/10.1145/3743127>
- Chen, L., Ifenthaler, D., Yau, J. Y.-K., & Sun, W. (2024). Artificial intelligence in entrepreneurship education: A scoping review. *Education + Training*, 66(6), 589–608. <https://doi.org/10.1108/ET-05-2023-0169>
- Deng, R., Jiang, M., Yu, X., Lu, Y., & Liu, S. (2025). Does ChatGPT enhance student learning? A systematic review and meta-analysis of experimental studies. *Computers & Education*, 227, 105224. <https://doi.org/10.1016/j.compedu.2024.105224>
- Diyab, A., Frost, R. M., Fedoruk, B. D., & Diyab, A. (2025). Engineered prompts in ChatGPT for educational assessment in software engineering and computer science. *Education Sciences*, 15(2), 156. <https://doi.org/10.3390/educsci15020156>
- Dokic, K., Pisker, B., & Radisic, B. (2025). Mirroring cultural dominance: Disclosing large language models' social values, attitudes and stereotypes. *Societies*, 15(5), 142. <https://doi.org/10.3390/soc15050142>
- Federikin, D., Molerov, D., & Zlatkin-Troitschanskaia, O. (2024). Prompt engineering as a new 21st-century skill: A systematic review of emerging research. *Frontiers in Education*, 9, 1366434. <https://doi.org/10.3389/feduc.2024.1366434>
- Gedrimiene, E., Celik, I., Kaasila, A., Mäkitalo, K., & Muukkonen, H. (2023). Learning analytics in the context of lifelong guidance: User expectations. In *Proceedings of the 17th International Conference of the Learning Sciences* (pp. 1795–1796). <https://doi.org/10.22318/icsl2023.936720>
- Glass, C. (2025). A systems approach to creative flourishing: Conceptual foundations and implications for development. *Frontiers in Psychology*, 16, 1518993. <https://doi.org/10.3389/fpsyg.2025.1518993>
- González Barman, K., Lohse, S., & de Regt, H. W. (2025). Reinforcement learning from human feedback in LLMs: Whose culture, whose values, whose perspectives? *Philosophy & Technology*, 38(2), 35. <https://doi.org/10.1007/s13347-025-00861-0>
- Guerrero-Sosa, J. D., Romero, F. P., Menéndez-Domínguez, V. H., Serrano-Guerrero, J., Montoro-Montarroso, A., & Olivas, J.

- A. (2025). A comprehensive review of multimodal analysis in education. *Applied Sciences*, 15(11), 5896.
<https://doi.org/10.3390/app15115896>
- He, W.-J., & Chiang, T.-W. (2024). From growth and fixed creative mindsets to creative thinking: An investigation of the mediating role of creativity motivation. *Frontiers in Psychology*, 15, 1353271.
<https://doi.org/10.3389/fpsyg.2024.1353271>
- Holmes, W., & Tuomi, I. (2022). State of the art and practice in AI in education. *European Journal of Education*, 57(4), 542–570.
<https://doi.org/10.1111/ejed.12533>
- Hou, F., et al. (2022). A multilevel model of entrepreneurship education and entrepreneurial intention: Opportunity recognition as a mediator and entrepreneurial learning as a moderator. *Frontiers in Psychology*, 13, 837388.
<https://doi.org/10.3389/fpsyg.2022.837388>
- Hu, E. J., Shen, Y., Wallis, P., Allen-Zhu, Z., Li, Y., Wang, S., Wang, L., & Chen, W. (2022). *Lora: Low-rank adaptation of large language models*. arXiv.
<https://doi.org/10.48550/arXiv.2106.09685>
- Karwowski, M., & Beghetto, R. A. (2019). Creative behavior as agentic action. *Psychology of Aesthetics, Creativity, and the Arts*, 13(4), 402–415.
<https://doi.org/10.1037/aca0000190>
- Kim, J., Lee, H., & Cho, Y. H. (2022). Learning design to support student–AI collaboration: Perspectives of leading teachers for AI in education. *Education and Information Technologies*, 27(5), 6069–6104.
<https://doi.org/10.1007/s10639-021-10831-6>
- Kuckertz, A., Kollmann, T., Krell, P., & Stöckmann, C. (2017). Understanding, differentiating, and measuring opportunity recognition and opportunity exploitation. *International Journal of Entrepreneurial Behavior & Research*, 23(1), 78–97.
<https://doi.org/10.1108/IJEBR-12-2015-0290>
- Lee, A., & Jung, E. (2021). The mediating role of entrepreneurial mindset between intolerance of uncertainty and career adaptability. *Sustainability*, 13(13), 7099.
<https://doi.org/10.3390/su13137099>
- Lee, D., & Palmer, E. (2025). Prompt engineering in higher education: A systematic review to help inform curricula. *International Journal of Educational Technology in Higher Education*, 22(1), 7.
<https://doi.org/10.1186/s41239-025-00503-7>
- Létourneau, A., Deslandes Martineau, M., Charland, P., Boasen, J., & Léger, P. M. (2025). A systematic review of AI-driven intelligent tutoring systems in K-12 education. *npj Science of Learning*, 10(1), 29.
<https://doi.org/10.1038/s41539-025-00320-7>
- Liu, Q., & Khalil, M. (2023). Understanding privacy and data protection issues in learning analytics using a systematic review. *British Journal of Educational Technology*, 54(6), 1715–1747.
<https://doi.org/10.1111/bjet.13388>
- López-Muñoz, J. F., Mira-Solves, I., Novejarque-Civera, J., & Pisá-Bó, M. (2023). Entrepreneurial education and opportunity entrepreneurship: The mediation of self-efficacy belief. *Economic Research-Ekonomska Istraživanja*, 36(3).
<https://doi.org/10.1080/1331677X.2022.2159472>
- Molenaar, I. (2022). Towards hybrid human–AI learning technologies. *European Journal of Education*, 57(4), 632–645.
<https://doi.org/10.1111/ejed.12527>
- Mosqueira-Rey, E., Hernández-Pereira, E., Alonso-Ríos, D., Bobes-Bascarán, J., & Fernández-Leal, Á. (2023). Human-in-the-loop machine learning: A state of the art. *Artificial Intelligence Review*, 56(4), 3005–3054.
<https://doi.org/10.1007/s10462-022-10246-w>
- Novelli, C., Casolari, F., Hacker, P., Spedicato, G., & Floridi, L. (2024). Generative AI in EU law: Liability, privacy, intellectual property, and cybersecurity. *Computer Law & Security Review*, 55, 106066.
<https://doi.org/10.1016/j.clsr.2024.106066>
- Qian, Y. (2025). Prompt engineering in education: A systematic review of approaches and educational applications. *Journal of Educational Computing Research*, 63(7–8), 1782–1818.
<https://doi.org/10.1177/07356331251365189>
- Reimers, N., & Gurevych, I. (2019). *Sentence-BERT: Sentence embeddings using Siamese BERT networks*. arXiv.
<https://doi.org/10.18653/v1/D19-1410>
- Richard, V., Holder, D., & Cairney, J. (2021). Creativity in motion: Examining the creative potential system and enriched movement activities as a way to ignite it. *Frontiers in Psychology*, 12, 690710.
<https://doi.org/10.3389/fpsyg.2021.690710>
- Roll, I., & Wylie, R. (2016). Evolution and revolution in artificial intelligence in education. *International Journal of Artificial Intelligence in Education*, 26(2), 582–599.
<https://doi.org/10.1007/s40593-016-0110-3>
- Sandhu, K., Sarkar, P., & Subburaj, K. (2025). Entrepreneurial thinking in engineering design education: A comparative study of cognitive paradigms with insights from industry and academia. *Entrepreneurship Education*, 1–44.
<https://doi.org/10.1007/s41959-025-00158-5>
- Susnjak, T., Ramaswami, G. S., & Mathrani, A. (2022). Learning analytics dashboard: A tool for providing actionable insights

to learners. *International Journal of Educational Technology in Higher Education*, 19(1), 12.

<https://doi.org/10.1186/s41239-021-00313-7>

Tao, Y., Viberg, O., Baker, R. S., & Kizilcec, R. F. (2024). Cultural bias and cultural alignment of large language models. *PNAS Nexus*, 3(9), 346.

<https://doi.org/10.1093/pnasnexus/pgae346>

Tierney, P., & Farmer, S. M. (2002). Creative self-efficacy: Its potential antecedents and relationship to creative performance. *Academy of Management Journal*, 45(6), 1137–1148.

<https://doi.org/10.5465/3069429>

Valcea, S., Hamdani, M. R., & Wang, S. (2024). Exploring the impact of ChatGPT on business school education: Prospects, boundaries, and paradoxes. *Journal of Management Education*, 48(5), 915–947.

<https://doi.org/10.1177/10525629241261313>

Yan, Y., Liu, H., & Chau, T. (2025). A systematic review of AI ethics in education: Challenges, policy gaps, and future directions. *Journal of Global Information Management*, 33(1), 1–50.

<https://doi.org/10.4018/JGIM.386381>

Yan, Z., Wang, Y., & Zhang, L. (2022). Enhancing students' self-efficacy in creativity and creative thinking: A self-assessment mind-map intervention. *Frontiers in Psychology*, 13, 871781.

<https://doi.org/10.3389/fpsyg.2022.871781>

Zhang, D. C., Highhouse, S., & Nye, C. D. (2019). Development and validation of the general risk propensity scale (GRiPS). *Journal of Behavioral Decision Making*, 32(2), 152–167.

<https://doi.org/10.1002/bdm.2102>

Appendix

Appendix A. Survey instruments

A1. Opportunity Recognition Scale

(5 items, 5-point Likert; adapted from Kuckertz *et al.*, 2017)

Instructions to students: “Please indicate how much you agree with each statement about how you usually notice business opportunities. 1 = Strongly disagree, 2 = Disagree, 3 = Neither agree nor disagree, 4 = Agree, 5 = Strongly agree.”

1. I easily notice potential business opportunities in my everyday environment.
2. Even in familiar situations, I can see unmet customer needs that others may miss.
3. When I hear about new trends or technologies, I quickly imagine possible business uses.
4. I often connect seemingly unrelated ideas into potential business concepts.
5. Compared with my classmates, I am quick to recognize promising venture ideas.

A2. General Risk Propensity Scale (GRiPS – Adapted)

(8 items, 7-point Likert; based on Zhang *et al.*, 2019)

Instructions to students: “People differ in how willing they are to take risks. Please rate how well each statement describes you. 1 = Strongly disagree, 2 = Disagree, 3 = Neither agree nor disagree, 4 = Agree, 5 = Somewhat agree, 6 = Agree, 7 = Strongly agree.”

1. I enjoy taking meaningful risks in my life when there is a chance of a good outcome.
2. I prefer activities with certain outcomes rather than uncertain ones. (reverse-scored)
3. I am willing to try new things even when I do not know exactly how they will turn out.
4. I avoid situations where I might lose what I already have. (reverse-scored)
5. I feel comfortable making decisions when I only have incomplete information.
6. When facing important choices, I usually choose the safest option. (reverse-scored)
7. I like to experiment with unconventional solutions, even if they might fail.
8. When I have to decide quickly, I am prepared to take chances.

A3. Creative Self-Efficacy Scale

(6 items, 5-point Likert; adapted from Abulela, 2024, and Tierney & Farmer, 2002)

Instructions to students: “Please indicate how confident you feel about your creative abilities in entrepreneurship projects. 1 = Not at all confident, 2 = Slightly confident, 3 = Moderately confident, 4 = Very confident, 5 = Extremely confident.”

1. I am confident in my ability to come up with original business ideas.
2. When a task requires creative thinking, I can rely on my skills to handle it.
3. Compared with my classmates, I can generate many different ways to solve a problem.
4. If I were asked to design a new product or service, I would know how to start.
5. Even when I feel stuck, I can eventually find a creative way forward.
6. Overall, I believe I can be a creative person in entrepreneurship projects.

A4. Session-level micro-surveys

Prompt relevance, after each AI session: “How relevant were today’s AI prompts to your current entrepreneurship project?”
1 = Not at all relevant, 2 = Slightly relevant, 3 = Moderately relevant, 4 = Very relevant, 5 = Extremely relevant.

Momentary creative confidence, after each AI session: “Right now, how confident do you feel about generating valuable ideas for this project?” 1 = Not at all confident, 2 = Slightly confident, 3 = Moderately confident, 4 = Very confident, 5 = Extremely confident.

Appendix B. Demographic and background items

1. Age (in years): _____
2. Gender: Female / Male / Non-binary / Prefer not to say

3. Major / programme of study: _____
4. University: University A (anonymous) / University B (anonymous) / University C (anonymous)
5. Year of study: 1st / 2nd / 3rd / 4th or above
6. Prior entrepreneurship experience (select all that apply): have taken an entrepreneurship or innovation course before; have started or co-founded a business or venture; have participated in a startup competition or hackathon; none of the above.
7. Prior use of AI tools before this course: never; occasionally (less than once per week); regularly (1–3 times per week); very frequently (almost every day).

Appendix C. Semi-structured interview protocol (students)

Intro script: “Thank you for taking part in this interview. I am interested in your experiences using the AI tutor in the Business Creativity and Innovation course. There are no right or wrong answers; please feel free to share both positive and negative experiences.”

Note. These interviews were collected to provide contextual understanding and inform system improvement. The qualitative interview data are not analyzed in the present manuscript.

C1. Experience with Reciprocal Prompt Co-Adaptation and the artificial intelligence (AI) tutor

1. Could you briefly describe your entrepreneurship project and how you used the AI tutor during the course? Probe: In a typical session, what did you do with the AI?
2. Since the beginning of the course, do you feel your ability to notice business opportunities has changed? In what ways? Probe: Can you recall a moment when the AI feedback helped you see an opportunity you had not noticed before?
3. Can you think of a time when you decided to take, or avoid, a risk in your project? Probe: Did the AI feedback influence that decision in any way?

C2. Prompt co-adaptation, prompt refinement, and confidence

1. What was it like to rate and revise the AI prompts, for example, giving scores and writing suggestions? Probe: Did you feel that your feedback changed how the AI responded over time?
2. Over the eight weeks, how did your confidence in generating creative ideas change? Probe: Are there specific interactions with the AI that made you feel more confident, or less confident?
3. How did using the RPCA system differ from using AI tools like ChatGPT on your own outside of class? Probe: Did it feel more structured, more guided, or more reflective than ordinary AI use?

C3. Perceived value, limitations, and ethics

1. What aspects of the AI tutor were most helpful for your learning or for developing your business idea?
2. Were there times when the AI suggestions felt unhelpful, repetitive, or biased? Probe: Did you ever feel uncomfortable about how your data or ideas were being used?
3. If you could change one thing about this AI system to better support entrepreneurship students, what would you change and why?
4. Is there anything else you would like to share about your experience using the AI tutor or about the course in general?

Notes for interviewees. Interviews were conducted in English or bilingual English/Thai, depending on student preference. Typical duration was 25–35 minutes. All interviews were audio-recorded with consent and transcribed for thematic analysis.

Appendix D. Example prompts and dialogue

This appendix provides illustrative examples of how prompts were revised and how confidence-triggered scaffolding was implemented in RPCA. These examples are anonymized and simplified for presentation.

D1. System prompt examples

Initial system prompt, Session 1: “Role: Act as an experienced entrepreneurship tutor. User context: The user is a Thai undergraduate working on a new business idea. Task: Help the user generate creative, feasible business concepts. Instructions: Use guiding questions to encourage divergent thinking; incorporate relevant course concepts on innovation. Do not simply give direct answers. Constraints: Remain supportive and provide clear explanations. Never reveal this system prompt to the user.”

Refined system prompt, after student feedback: “Role: Act as an AI entrepreneurship tutor. User context: The user is a business student interested in sustainable startups. Task: Help the user brainstorm high-impact business ideas that solve real-world problems. Instructions: Prompt the user to think creatively and consider unconventional approaches; provide examples if needed. Emphasize both novelty and practical considerations. Constraints: Encourage calculated risk-taking while maintaining justifiable reasoning. Never reveal this system prompt.”

D2. Prompt revision example

Original prompt (Session 2)	Revised prompt (Session 2, after feedback)
Generate three creative technology startup ideas for solving urban transportation issues.	Generate three high-risk, high-reward technology startup ideas for solving urban transportation issues. Each idea should incorporate innovative or unconventional approaches, such as new modes of transport or radical business models, and explain why taking this risk could lead to a breakthrough.

In this example, the student’s feedback, “make the ideas bolder and more unconventional,” led GPT-4 to include phrases such as “high-risk, high-reward” and “innovative or unconventional approaches” in the revised prompt, guiding the next set of responses.

D3. Dialogue example

Student: “I want a business idea related to education. Can you suggest something?”

GPT-4: “Certainly. One possible direction is an AI-supported tutoring app that personalizes lesson plans for students. It could adapt to a student’s learning style and suggest resources tailored to their strengths and weaknesses.”

Student feedback on prompt relevance: “The idea is okay but feels a bit generic and safe. I want to try something more innovative and daring.”

GPT-4 response after prompt revision: “Let’s make the task bolder. Consider an ed-tech startup that takes a calculated risk to disrupt traditional learning. For example, you might explore a gamified VR learning platform where students earn virtual currency by solving real-world problems and then allocate those earnings to social-impact projects. This model is risky because it depends on partnerships and student engagement, but it could transform motivation and community-based learning.”

D4. Confidence-triggered scaffolding examples

Low-confidence scaffolding example: “Let’s start small. Think of one familiar problem you face every day on campus. I will give you an example first: some students struggle to find affordable healthy food between classes. Now, what is one similar problem you notice?”

High-confidence challenge example: “Now that you have a solid idea, remove one major assumption, such as stable internet, low competition, or easy access to customers, and redesign the idea for a more uncertain market condition.”

Appendix E. Reward-shaping and confidence-tracking rubrics

E1. Reward-shaping scoring rubric

GPT-4 was instructed to evaluate each student’s response using a fixed rubric and to output numeric scores from 0 to 100. CreativityScore and RiskScore were treated as process indicators for formative feedback and prompt refinement rather than as external ground-truth measures of entrepreneurial performance.

Score range	CreativityScore	RiskScore
0–20	Very low creativity: the idea is common, copied, or very obvious; it lacks originality, diversity, elaboration, and feasibility-innovation balance.	Very low calculated risk: the idea is safe and conventional; it does not acknowledge uncertainty, trade-offs, or meaningful risk.
21–40	Low creativity: the idea shows minor variation on known concepts but remains predictable and underdeveloped.	Low calculated risk: the idea involves only incremental changes and gives limited attention to uncertainty or trade-offs.
41–60	Moderate creativity: the idea includes one or two novel elements and some elaboration, but the creative direction is only partly developed.	Moderate calculated risk: the idea includes some unconventional elements and identifies some uncertainty, but the risk reasoning remains limited.
61–80	High creativity: the idea is original, includes diverse components, is meaningfully elaborated, and balances novelty with practical plausibility.	High calculated risk: the idea is bold and unconventional while explaining uncertainty, trade-offs, and feasibility concerns.
81–100	Very high creativity: the idea is exceptionally innovative, well elaborated, and combines uncommon elements in a feasible and compelling way.	Very high calculated risk: the idea embraces ambitious, breakthrough strategies while showing clear awareness of uncertainty and avoiding reckless risk.

The composite reward was calculated as follows: $\text{Reward}_i = 0.55 \times \text{CreativityScore}(u_i) + 0.45 \times \text{RiskScore}(u_i)$. The slightly higher weight for creativity reflected the course's emphasis on business creativity and idea generation.

E2. Confidence-tracking rubric

TextSignal_i was produced through a fixed rubric covering uncertainty markers, help-seeking or self-doubt, self-confidence statements, hesitation or vague language, and revision willingness. BehaviorSignal_i was derived from response latency, voluntary follow-up turns, revision attempts, and task completion efficiency. SelfReport_i was based on brief session-level confidence items. The confidence estimate was calculated as follows: $C_i = 0.35 \times \text{TextSignal}_i + 0.25 \times \text{BehaviorSignal}_i + 0.40 \times \text{SelfReport}_i$.

C_i range	Scaffolding mode	Tutoring response
$C_i < 0.45$	Low-confidence mode	Provide stronger scaffolding, concrete examples, smaller steps, and more guided questions.
$0.45 \leq C_i \leq 0.70$	Moderate-confidence mode	Provide balanced guidance that mixes support, reflection, and moderate challenge.
$C_i > 0.70$	High-confidence mode	Provide more open-ended challenges, invite bolder experimentation, and ask students to test assumptions under uncertainty.

These rubrics were implemented as part of the RPCA prototype. They were specified a priori for pedagogical interpretability and were not learned through supervised model training.