

提示工程视角下GAI促进大学生批判性思维能力的研究

陈璐, 吴嘉琳, 秦炜炜*

苏州大学教育学院, 江苏 苏州

收稿日期: 2025年7月9日; 录用日期: 2025年8月18日; 发布日期: 2025年8月27日

摘要

随着生成式人工智能(GAI)在高等教育中的快速普及和深度应用, 其对大学生批判性思维能力的影响引发广泛关注。本研究立足提示工程视角, 通过对50位不同专业、年级本科生的半结构化访谈, 结合MAXQDA2022软件进行数据分析, 系统揭示了GAI交互对大学生批判性思维能力的影响机制。研究发现, 提示设计的动态演进、过程调节的深度赋能及元认知能力的迭代构成了影响大学生批判性思维能力的关健路径, 揭示了提示工程如何将GAI从工具转化为批判性思维发展的有效载体。研究在理论上丰富了GAI与批判性思维关系的研究维度, 拓展了提示工程在教育场景的理论内涵; 在实践中为高校利用GAI培养批判性思维提供了现实路径。未来可进一步探索提示工程在多场景中的应用, 为智能化时代大学生认知能力发展提供有效支撑。

关键词

生成式人工智能, 提示工程, 批判性思维, 大学生, 影响机制

Research on the Influence Mechanism of GAI Promoting Critical Thinking Skills of College Students from the Perspective of Prompt Engineering

Lu Chen, Jialin Wu, Weiwei Qin*

School of Education, Soochow University, Suzhou Jiangsu

Received: Jul. 9th, 2025; accepted: Aug. 18th, 2025; published: Aug. 27th, 2025

*通讯作者。

文章引用: 陈璐, 吴嘉琳, 秦炜炜. 提示工程视角下 GAI 促进大学生批判性思维能力的影响机制研究[J]. 创新教育研究, 2025, 13(8): 568-581. DOI: 10.12677/ces.2025.138634

Abstract

With the rapid popularization and in-depth application of generative artificial intelligence (GAI) in higher education, its impact on college students' critical thinking ability has attracted widespread attention. Based on the perspective of prompt engineering, this study systematically reveals the influence mechanism of GAI interaction on college students' critical thinking ability through semi-structured interviews with 50 undergraduates of different majors and grades, combined with data analysis by MAXQDA2022 software. The study finds that the dynamic evolution of prompt design, the deep empowerment of process modulation, and the iteration of metacognitive ability constitute the key paths that affect college students' critical thinking ability, revealing how prompt engineering can transform GAI from a tool to an effective carrier for critical thinking development. Theoretically, this study enriches the research dimension of the relationship between GAI and critical thinking, and expands the theoretical connotation of prompt engineering in the educational scene. In practice, it provides a realistic way for colleges and universities to use GAI to cultivate critical thinking. In the future, the application of prompt engineering in multiple scenarios can be further explored to provide effective support for the cognitive development of college students in the intelligent era.

Keywords

Generative Artificial Intelligence, Prompt Engineering, Critical Thinking, College Students, Mechanisms of Influence

Copyright © 2025 by author(s) and Hans Publishers Inc.

This work is licensed under the Creative Commons Attribution International License (CC BY 4.0).

<http://creativecommons.org/licenses/by/4.0/>



Open Access

1. 引言

随着生成式人工智能(Generative Artificial Intelligence, GAI)技术的突破性发展,以大型语言模型(Large Language Model, LLM)为代表的新型智能工具正深度融入高等教育体系。以 ChatGPT、DeepSeek、Claude 等为代表的 GAI 平台,凭借其强大的自然语言理解和生成能力,为大学生提供了即时知识解答、文本生成、问题分析和创意激发服务,深刻挑战并重塑着传统的教与学范式。然而, GAI 在高等教育中的普及浪潮也引发学界的忧虑。有学者提出,在人工智能加速渗透的时代背景下,高等教育的核心使命应当聚焦于学生的批判性思维、道德判断力、创新能力等人类独有的品质[1]。现有研究表面,大学生倾向于被动接受 GAI 输出的“现成答案”,将其视为“知识黑箱”,这种浅层化的人机交互模式,极易诱发一系列认知惰性問題,可能导致学生批判性思维(Critical Thinking)能力的弱化[2]。究其本质,这源于当前许多人机交互停留于表层的“提问-获取答案”过程,未能有效激发学生进行深度的认知建构和元认知反思。

值得注意的是, GAI 的输出质量和认知深度并非固定不变,而是高度依赖用户输入的提示(Prompt)的质量与结构。提示工程(Prompt Engineering)作为优化人机交互的核心技术,旨在通过精心设计结构化、策略性、引导性的指令,主动引导 LLM 进行更深入的推理、多视角的分析、严格的证据评估、假设验证以及清晰的解释生成[3]。本质上,提示工程通过设计特定的认知任务框架,将用户置于主动引导和深度参与的位置,而非被动接收端。因此,在 GAI 深度融入高等教育的背景下,如何有效利用提示工程这一杠杆,将潜在的思维弱化风险转化为批判性思维能力发展的新契机,成为一个亟待探索的关键问题。当前研究多聚焦 GAI 风险管控或学科应用,而从提示工程视角系统探究其如何影响并塑造大学生批判性思维发展机制的研究尚显不足。据此,本研究提出如下核心问题:在高等教育场景中,提示工程驱动下的 GAI

交互如何具体影响大学生批判性思维能力？其作用机制为何？

2. 文献综述

2.1. GAI 在高等教育中的应用现状

GAI 已深度融入高等教育的教、学与评核心环节。在教学辅助领域，GAI 通过自动化生成课程材料、定制个性化学习路径及提供实时答疑，显著提升教学效率[4]；在学术实践中，学生广泛利用其完成文献综述框架构建、数据可视化代码生成及论文语言润色[5]；评估体系则逐步引入 GAI 驱动的自动化评分工具[6]；在协作学习中，GAI 能够有效促进学生创造性潜能发展[7]。然而，当前 GAI 呈现出明显的工具化浅层应用倾向，学生倾向于被动采纳 GAI 输出的结论性内容，如生成论文论点、获得解题答案等，将其视为“权威知识源”而非激发深度思考的起点[8]，导致人机交互多停留于事实查询、文本改写等基础问答层面，缺乏引导多轮质疑与逻辑验证的深度设计。实证研究进一步指出，目前已有不少大学生意识到 GAI 的优势与不足，在利用其开展学习时能主动、负责任地与 GAI 交互[9]，但仍有大学生存在过度依赖 GAI 的现象，导致在开放式问题解决中表现出论证严谨性较低与原创性不高，这一现状凸显了重构人机协作范式、将 GAI 从“答案供给者”转型为“思维脚手架”的紧迫性。

2.2. 批判性思维

批判性思维作为高等教育的核心素养，本质是基于证据主动分析、评估与重构认知的元能力，涵盖解释、分析、推论、评估、说明及自我调节六大核心维度[10]。传统教学模式依赖案例研讨、苏格拉底问答等训练该能力，但 GAI 的普及带来多重结构性挑战：其一，信息甄别意识钝化，流利性错觉(Fluency Illusion)使学生易被 GAI 生成逻辑连贯但潜在错误的文本误导，缺乏对 GAI 生成内容的准确性、偏见和来源可靠性的主动质疑和核查[11]；其二，逻辑推理能力下降，满足于 GAI 提供的结论，忽视其论证过程的严密性、证据的充分性和潜在的逻辑谬误；其三，创新思维受限，认知卸载(Cognitive Offloading)诱使学生将复杂推理任务移交 GAI，弱化了自身进行复杂问题拆解、多角度深度分析和原创性观点构建的能力[12]；其四，批判对象认知错位，GAI 的“算法黑箱”特性使大学生因缺乏算法透明性而陷入批判失焦，如将判断焦点从生成内容的可信度转向生成过程的可靠性[13]。这些挑战表明，亟需构建适配 GAI 环境的批判性思维训练框架。

2.3. 提示工程

提示工程(Prompt Engineering)是通过结构化指令优化生成式人工智能(GAI)输出的关键技术，其核心是将人类意图转化为机器可执行的指令，解决 GAI 输出的模糊性、偏见及伦理风险问题。在教育领域，它已超越技术操作，成为激发深度认知协作的思维训练工具[14]。其核心方法，如 CRISPE 框架、“Think step by step”元提示以及角色重构提示[15]，通过迫使用户在交互前厘清问题、预设标准、拆解论证并评估输出，直接训练分析、评估与反思等批判性思维核心技能。若有效加以利用，此类设计能将学生从答案“被动接收者”转化为“主动协作者”，通过创编提示需主动界定问题、构建假设这一认知主动性路径和针对性提示强制质疑、验证与意义重构这一深度加工路径，在动态交互中激活并锻炼其批判性思维的核心技能。

3. 研究设计

3.1. 研究方法

本研究采用半结构化访谈法作为质性研究核心方法，此方法在预设核心问题框架的基础上，允许研

研究者根据受访者回答灵活追问,从而有效挖掘参与者复杂具体的主观体验和认知过程。访谈时间控制在15~20分钟之间,访谈提纲设计紧密围绕核心研究问题展开,主要涵盖以下维度:

- (1) GAI 使用习惯与场景(频率、目的、常用任务类型);
- (2) 提示工程驱动下的 GAI 交互具体过程(如何设计提示引导分析、推理、评估、反思等);
- (3) 感知到的 GAI 交互对自身思考方式的影响;
- (4) 遇到的挑战与局限性(如 GAI 输出的可靠性判断、认知负荷等)。

3.2. 数据来源

本研究采用目的性抽样结合最大差异抽样策略,最大差异抽样的核心标准在于捕捉与研究主题最相关的关键差异维度,具体选定为“年级”和“专业”。受访者选择标准主要包括:具备丰富的 GAI 交互经验;能够清晰表达自己的观点和看法。研究通过对 50 位不同专业、年级的本科生进行半结构化访谈,一大二大三受访学生比例为 1:1.25:4,人文社科类占 52%,理工科类占 48%。尽管抽样策略力求在年级和专业维度上实现最大差异,样本仍存在倾斜性局限,大三学生占比显著偏高,可能使研究结果更多反映高年级视角。访谈前均严格获得参与者的知情同意,明确告知研究目的并承诺数据保密与匿名处理原则。访谈过程在征得同意后录音,收集的访谈录音资料通过“讯飞听见”软件转录为逐字稿,并由人工校对保证准确性,形成 50 份原始文本数据集(一人一份),共计 92,340 字。为保护隐私,所有访谈资料均进行匿名化处理,并对接受访谈者按顺序进行编号 A1、A2、A3……以此列推。

3.3. 数据分析及模型构建

MAXQDA 作为一款专业的质性研究分析软件,能高效支持研究者对访谈、观察记录等原始文本进行系统化编码、范畴梳理与关系挖掘,在文本数据处理与质性分析中被广泛应用。本研究借助 MAXQDA2022 软件对原始文本数据进行分析整理。

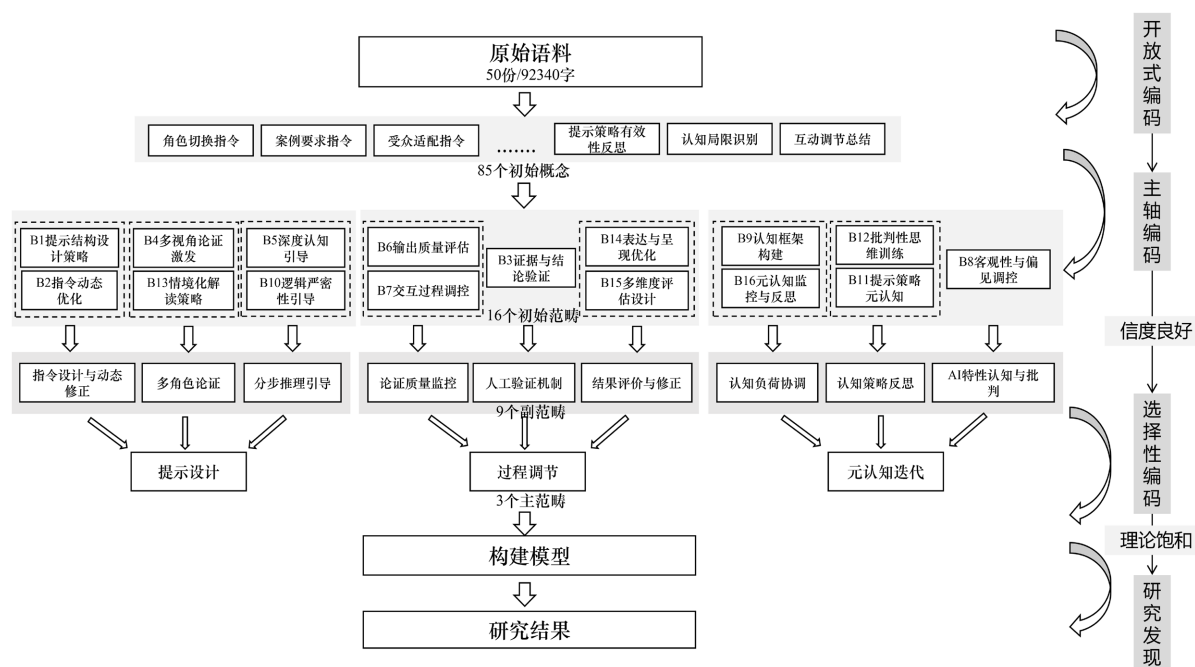


Figure 1. Interview data encoding process

图 1. 访谈数据编码流程

为确保编码过程的客观性，本研究采用两位编码员独立编码的方式进行信度检验。随机抽取 30% 的原始语料，由两位具备质性研究经验的编码员按照既定编码框架进行独立编码。采用 Cohen's Kappa 系数计算编码一致性，结果显示：开放式编码阶段 Kappa = 0.82，主轴编码阶段 Kappa = 0.78，选择性编码阶段 Kappa = 0.85，说明本研究编码过程具有较高的信度。对于编码不一致的案例，通过两位编码员与研究组长共同讨论达成共识，进一步修正编码框架，具体编码过程如图 1 所示。

3.3.1. 开放式编码

通过对访谈资料的逐句阅读、概念命名、编码整合，共提炼出如多轮追问、修正指令、结果评价等 85 个初始概念，在此基础上进一步将存在交叉的初始概念归纳压缩为 16 个初始范畴，如表 1 所示。

Table 1. Example of open coding analytics
表 1. 开放式编码分析举例

原始语句	初始概念	初始范畴
A18: 有时候我会给 GAI 一个角色设定，你现在是团支书，要策划团日活动	角色切换指令	B1 提示结构设计策略
A2: 提问时我会写上要求 GAI 用 3 个具体案例支撑观点，避免空谈理论	案例要求指令	
A30: 写教案时让 GAI 解释加密解密时，我会说明用适合高中生的语言	受众适配指令	
A41: 我会让 GAI 用学术论文的格式呈现观点，包括摘要、关键词和参考文献	结构化对比指令	
A50: 要求 GAI 用思维导图呈现思政课上某个理论的因果关系	可视化要求	
A16: 让 GAI 帮我整理好所需要的引用格式	格式规范指令	B2 指令动态优化
A7: 发现 GAI 答非所问后，我会追加请围绕具体时间界定等等	范围限定修正	
A19: 学习 JAVA 时遇到语法结构与课本冲突，最后放弃了 GAI 所提供的内容	术语纠错指令	
A23: GAI 梳理文献的时候容易出现大量重复语句，需要我补充新的分析角度	冗余修正指令	
A1: 分析市场趋势时，我会提醒 GAI 补充被忽略的变量影响	变量遗漏修正	
A31: 通过不断的修正指令，生成的文案从华丽和浮夸变得越来越切实和具体	具体化修正	B3 证据与结论验证
A12: 如果给出的答案还有一些我没有学的知识，我会一层一层继续问下去	语境补充修正	
A13: 我会要求标注文献来源，并且我会人工根据提供的来源精读核心文献	证据溯源要求	
A28: 看它网络上给出的一些解释是什么样的，然后再与 AI 的作对比	证据链对比	
A44: 我会让 GAI 先给出结论，再反过来请你推翻这个结论，看看能否成立	反例要求	
A16: 我进行真正的实验探究，来确定是否这些溶剂哪些是真正有效的	假设验证引导	B4 多视角论证激发
A5: 我会要求他标注文献的来源	可验证性标注	
A18: 仔细查了它给的文献，发现知网和各个地方都找不到，是它编造出来的	数据真实性核查	
A20: 让 GAI 从社会学和计算机科学两个角度分析算法偏见	跨学科视角引导	
A38: GAI 由于是集合互联网的数据，在无形中会存在严重的种族偏见	伦理框架切换	
A22: 打辩论的时候我觉得很好用，可以让 GAI 模拟反方立场	对立视角预设	B5 深度认知引导
A10: 国产生成式人工智能底层逻辑从中国的角度出发的，思维方式与中国人有着类似的方面，而 ChatGPT 等的底层逻辑是从西方文化的角度出发	文化视角引导	
A12: 我会让它们帮我把每一步的任务要做什么给我列出来	分步拆解指令	
A28: 讨论某个学术观点时，我会让 GAI 先假设该观点不成立，再推导可能产生的矛盾，以此验证其合理性	反证法引导	
A44: 在创造性和深度分析方面仍有局限	复杂性强调	
A26: 因为我自己偷懒不想学习计算机的课程，所以都依赖 GAI 工具帮我生成课程作业，导致期末考试的时候得到了一个很低的绩点	因果链追溯	

续表

A37: 让 GAI 点评课程论文时, 我会要求它多方面分别点评, 确保点评有条理	结构化点评要求	B6 输出质量评估
A35: 可用但需思考判断	可信度标注	
A44: 偏理工科的领域 AI 判断大部分正确, 但偏文科的内容解决起来比较困难	方法适用性评估	
A33: 所给的观点太少, 并且缺乏深度理解	论证强度评估	
A42: 更多反思自身判断和验证不足	错误路径提示	
A45: 有时输出的策划都是没办法实现的套路化的模板	创新性评估	B7 交互过程调控
A36: 每次提问后, 我都会让 GAI 先重复下我的问题核心, 确认它懂了	问题确认指令	
A37: 我询问一些关于历史的可能涉及敏感问题 Deepseek 并没有回答我	回避点质疑	
A38: 发现 GAI 推理跳步了, 我会让它把中间的推导过程补上	过程补全	
A39: 讨论跑题时, 我会提醒 GAI 回到大学生就业难的核心原因上来	焦点回归	
A40: 让 GAI 设计课堂活动时, 我会让它标注哪些环节需要学生分组讨论	互动设计	B8 客观性与偏见调控
A41: 我会跟 GAI 说这部分详细点说, 下部分简单概括下就行	节奏控制指令	
A15: GAI 由于是集合互联网的数据, 在无形中会存在严重的种族偏见	偏见识别引导	
A32: 一些基本的伦理或者逻辑错误我能自己判断出来, 如果是错的直接抛弃	客观性修正	
A44: 看到 GAI 说所有学生都适合在线教育, 我会让它加上些限定条件	绝对化修正	
A37: 以前遇到不会的问题都是通过夸克拍题来解决, 现在直接问 GAI 就可以	信息源限定	B9 认知框架构建
A6: 我希望能给我正确合理且符合现实社会的答案	表述要求	
A11: 了解 GAI 的发展历程和一些热门的生成式人工智能模型	知识要素提取	
A48: 我会让 GAI 先总结下 SWOT 分析法的框架, 再用到这个项目评估里	框架迁移引导	
A49: 讨论后现代主义文学时, 我会让 GAI 分析它和现代主义的继承与革新	传承关系分析	
A26: GAI 依赖人们不断的输入和数据的训练变得越来越智能, 而人们则在依赖 GAI 工具中人的能力越来越退步	依赖关系标注	B10 逻辑严密性引导
A47: 有一次小组作业需要设计一个罕见病科普海报, 我用 GAI 生成了患者病程的时间轴和通俗版病理机制图	跨学科整合	
A11: 分析学习技术时, 我会让 GAI 先明确这个概念的核心要素和应用边界	概念界定要求	
A33: 所讲述的内容没有逻辑上的联系	逻辑断层标记	
A44: 我会质疑这个结论的前提是否忽略了学科差异和学习方法的多样性	前提合理性质疑	
A21: 到 GAI 用“心理异常”这个词时, 我会让它具体说明是指临床诊断标准还是日常情绪波动, 避免概念混淆	概念模糊质疑	B11 提示策略元认知
A33: 我会要求他标注文献的来源	依据标注要求	
A29: 我会记录 GAI 的回答质量差异, 发现分点方式更易得到精准结果	提示效果记录	
A31: 我大概 2/3 都归咎于自己没有事先验证, 其他我会反思是 GAI 的局限性	失败原因溯源	
A27: 当 GAI 输出偏离预期的时候, 我第一反应会修正指令, 但是有时候修正指令还是没用, 这时候我倾向于自己完成	提示迭代对比	
A29: 长文本生成时 GAI 易重复, 短文本更精准, 这影响我对输入长度的控制	长度效应分析	B12 批判性思维训练
A44: 偏理工科我认为 AI 的判断大部分正确, 但偏文科 AI 解决起来比较困难	复杂度效应分析	
A19: 让 GAI 写辩论稿时, 我会要求它提前预判对方可能的反驳并准备回应	预判反驳设计	
A4: 使用 GAI 前我会预判它在复杂公式推导中可能出现步骤遗漏, 因此会提前要求它详细列出每一步演算过程	薄弱点预判	
A30: 辩证看待, 会学习的人会合理运用该工具	辩证性思维引导	
A42: 用 GAI 做项目时, 我会让它先分析可能出现的风险及应对措施	防御性思维引导	

续表

A16: 分析教育政策时, 我会让 GAI 关联当时的社会经济背景来解读	背景关联指令	B13 情境化 解读策略
A23: 在情感大于技能要求方面, 比如前些天我妈妈生日, 在给她订制的电子贺卡上可以选择祝福语, 有自动生成的, 也可以自行填写	情境差异分析	
A21: 让 GAI 设计教学方案时, 我会要求它考虑不同年龄段学生的接受差异	群体差异分析	
A23: 让 GAI 帮我对改革开放以来我国教育管理学的发展历程进行分阶段	历史脉络梳理	
A16: 我会让 GAI 结合改革开放后中外交流增多的时代背景来解读其出台原因	时代背景关联	
A10: 让 GAI 生成活动策划案时, 我会特别说明概括核心流程, 避免内容冗长	精简要求指令	B14 表达与 呈现优化
A34: 让 GAI 评估学习效果时, 我会要求它用具体指标如知识点掌握率来呈现	指标化要求	
A47: 生成了患者病程的时间轴和通俗版病理机制图	生活化转化要求	
A36: 要求 DeepSeek 基于 smart 方法生成一个表格, 并不断补充信息完善表格	结构化输出	
A39: 生成一个流程图, 直观详细	具象化引导	
A23: 对发展历程进行分阶段	时间维度划分	B15 多维度 评估设计
A31: 让 GAI 计算实验数据时, 我会要求它分析可能的误差来源及范围	误差分析要求	
A26: GAI 生成代码时经常忽略版本差异, 结合具体开发环境方面存在局限	方法局限说明	
A21: GAI 的风险肯定是有的, 比如说这个侵权的风险	风险预判要求	
A45: 用 GAI 写完论文后, 我会让它检查自身论述中的逻辑漏洞和论据不足	自我漏洞识别	
A26: 但是通过不断具象化我的指令, 我发现我可以不通过 GAI 来生成结果, 因为我已经知道自己想要什么样的结果了	提示策略有效性 反思	B16 元认知 监控与反思
A50: 后来发现不同的 GAI 工具有各自丰富的优势和功能	认知局限识别	
A27: 每次用 GAI 完成任务后, 我会总结哪些指令调节方式最有效, 比如分点提问比笼统提问更好	互动调节总结	

3.3.2. 主轴编码

本研究进一步对开放式编码过程获得的 16 个初始范畴分析范畴间关系, 归纳整理指令设计与动态修正、多角色论证、分步推理引导、论证质量监控、人工验证机制、结果评价与修正、认知负荷协调、认知策略反思、AI 特性认知与批判 9 个副范畴, 最终形成 3 个主范畴, 分别为提示设计、过程调节、元认知迭代, 如表 2 所示。

Table 2. Main categories formed by axial coding and their connotations
表 2. 主轴式编码形成主范畴及范畴内涵

范畴内涵	副范畴	主范畴
构造有效的提示词来引导 GAI 生成期望的内容, 并在交互过程中根据进展实时调整和优化提示	指令设计与动态修正	提示设计
设计提示或互动方式, 促使 GAI 模拟不同立场、角色或视角参与论证	多角色论证	
将复杂的论证问题分解为更小的、逻辑连贯的步骤, 并通过提示引导 GAI 一步步进行推理	分步推理引导	
在论证过程中或阶段结束时, 有意识地评估 GAI 输出的质量	论证质量监控	过程调节
在关键步骤处, 引入人工判断、事实核查, 以弥补可能存在的错误、偏见或知识盲区	人工验证机制	
对论证的最终结论或产出进行系统性评价, 并根据评价结果进行必要的完善	结果评价与修正	
有意识地管理自身的认知资源分配, 识别何时任务过于复杂导致负荷过载, 何时负荷不足, 并做出相应调整	认知负荷协调	元认知迭代
对自身使用的方法、思维过程进行回顾、评价和优化的机制	认知策略反思	
对 GAI 的能力优势、固有局限保持清醒认识和批判态度	AI 特性认知与批判	

3.3.3. 选择性编码

通过对主范畴之间关系的反复考察和分析,将核心范畴确定为“GAI 促进大学生批判性思维能力的影响机制”,据此,本研究构建了“GAI 促进大学生批判性思维能力的影响机制模型”,如图 2 所示。

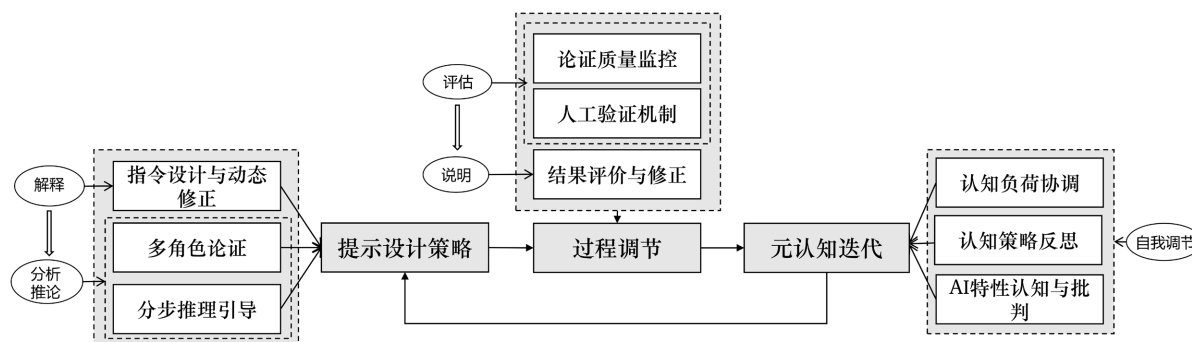


Figure 2. A model of the influence mechanism of GAI in promoting college students' critical thinking ability

图 2. GAI 促进大学生批判性思维能力的影响机制模型

本研究的核心范畴为“GAI 促进大学生批判性思维能力的影响机制”,围绕该核心范畴的“故事线”可以概括为:大学生在与 GAI 的深度交互中,将 GAI 从单向信息供给工具转化为多维度思维训练载体,系统性激活并强化批判性思维的核心能力。其一,提示设计阶段。问题界定时首先依托提示设计策略,通过指令设计与动态修正、多角色论证、分步推理引导,将隐性思维需求外化为可操作的指令要求,驱动解释、分析、推论等批判性思维基础活动的开展;其二,过程调节阶段。借助论证质量监控、人工验证机制实施评估,迫使学生持续审视生成内容的逻辑严密性、证据充分性与潜在偏见,从而直接训练推论能力与评估能力,动态调控交互流程;其三,元认知迭代阶段。认知负荷协调、认知策略反思、AI 特性认知与批判帮助大学生超越表层的答案接收,将多角度分析内容转化为个人化认知产物。学生对生成结果的批判性反馈及提示迭代修正行为,自然触发自我调节能力的发展,优化提示设计的精准性,强化过程调节的有效性,完成从工具使用到元认知能力迭代的跃迁,三者循环互动,逐步建构并提升大学生的批判性思维能力。

3.3.4. 理论饱和度和检验

在完成对前 50 份访谈资料的分析并初步构建模型后,继续分析预留的最后 5 份访谈资料。分析结果显示,这 5 份资料中出现的新信息、新观点均能被模型中已有的范畴充分解释和覆盖,没有产生新的重要概念、范畴或关系。因此,判定该模型已达到理论饱和,具有较好的解释力。

该模型在理论上呼应了自我调节学习模型等经典理论的核心思想,即学习者通过循环往复的主动监控与调整实现能力提升。本模型在 GAI 特定情境下,具体化了这一循环互动,揭示了 GAI 环境中自我调节对象与机制的根本性拓展,即大学生既要调节自身认知过程,如监控论证质量、反思认知策略,更要通过精心的提示设计动态调控 GAI 的思维输出方向与质量,将 GAI 转化为可调控的思维训练伙伴。此外,模型强调了基于对 GAI 特性认知的批判性反馈成为触发元认知迭代的关键动力,突显在智能技术深度融入学习的背景下,对技术中介本质的批判性理解已成为元认知发展不可或缺的新要素。因此,本模型既丰富了技术环境下批判性思维培养的理论视角,也为实践中优化 GAI 辅助教学策略提供了可操作的路径指引。

4. 研究发现

4.1. 提示设计的动态演进

研究发现,大学生与 GAI 的交互呈现阶梯式能力跃迁,逐步从模糊指令转向系统性提示设计,驱动自身主动完成解释、分析、推论的过程。初始阶段,宽泛提问往往引发 GAI 输出泛化、逻辑分散的结果,当输出内容偏离预期时,受访者 A10 则表示“我第一反应是质疑自己的指令然后进一步去完善,表达或者描述的更详细一点,更符合我的要求”,如 A5 最初以“分析短视频对社交的影响”提问,因输出空泛,通过限定分析框架,将提示词细化为“从青少年社交模式变化、隐私风险两个角度结合具体案例分析”,通过指令设计与动态修正,迫使 GAI 输出更聚焦的回答,揭示了大学生从归咎工具到审视自身表达的认知转向。

随着交互经验积累,提示设计向多角色论证与分步推理引导纵深发展。在多角色论证层面,A18 表示“有时候我会给 GAI 一个角色设定,就说你现在就是一名团支书,要策划团日活动或者写学期活动总结等等”,将抽象任务锚定到具体角色的权责、场景与受众。这种情境化指令不仅提升 GAI 输出的有效性,更驱动大学生在设定角色时预先考虑其利益关联、认知局限和表达风格,客观上训练了立场分析与论证建构能力。

在分步推理引导层面,大学生通过问题拆解、阶段设问、递进引导降低思维复杂度。认知负荷理论指出,工作记忆的有限容量是深度学习的核心瓶颈,大量概念、关系同时处理,极易超出工作记忆容量。因此,将庞大探索任务拆解为有序子目标有助于优化认知资源分配。如 A27“在了解没有涉足过的领域时,我会先让 GAI 告诉我有哪几个方向可以去了解,再从它提供的信息里继续提问”,将 GAI 转化为思维脚手架,每个步骤仅需处理有限信息,避免一次性涌入所有关联概念导致的超载,又促使自身在接收、筛选、拓展的循环中强化信息评估与路径规划能力。这种模块化提示设计让 GAI 输出更具逻辑层次,也迫使使用者在指令设计时预先完成思维结构化,明确各阶段核心任务,从而在交互前就启动解释、分析、推论的思维链路。

这种动态的提示设计修正行为标志着大学生从被动接受转向主动驾驭工具,通过指令的精准调控,逐步唤醒分析、论证等批判性思维的基础能力,实现从“问结果”到“控过程”的思维觉醒。

4.2. 过程调节的深度赋能

交互过程中,大学生围绕“评估-说明”的逻辑,通过论证质量核验、人工验证机制、结果动态修正的协同联动,形成过程调节,推动批判性思维从浅层质疑向深度建构升级。在论证质量核验上,大学生既关注逻辑自洽性,也重视实证可靠性,A20 曾“碰见给出的回答和已有理论现实矛盾的情况”,A31 指出“GAI 写的文本缺少文学素养和情感深度,有时候语言很人机味缺少逻辑性”;A19 用 GAI 制作思政课 PPT 时要求 GAI 给出一份大纲并制作 PPT,但“GAI 制作的 PPT 模板内容字体和排版都不合适”。大学生对 GAI 输出内容的评估已突破了单纯的对错判断,上升为对回答质量更精细的研判,从内容层面的逻辑严谨性、学术素养校验,到形式层面的呈现适配性考量,大学生通过主动拆解质量要素、重构评估标准,将自身认知需求转化为可操作的评估指标,既实现了对信息质量的严格把控,也强化了人与智能工具协同中的主体性与创造性。

在人工验证机制上,大学生突破对技术的盲从,构建多渠道交叉核验的习惯。A23 在撰写文献综述时借助 GAI 但“去知网搜了标题,根本没有它说的论文,后来发现它自己编了文献”,面对 GAI 提及的未知来源信息,大学生并未直接信任,而是选择去权威网站检索核对;A42 表示“GAI 会随机编写数据”,发现数据可疑时,会同时比对行业报告、政府公开数据等,确认“数据确实存在才敢用”。大学生在与

GAI 交互的过程中不轻易接受工具输出的“确定性结论”，而是通过时效性核查、多渠道比对等方式，把被动接收转化为主动验证，强化了对信息的甄别与评估能力。

在结果评价与修正上，大学生形成发现问题后反馈具体情况、优化指令后再获取更精准输出的循环。A7 使用 GAI 写代码时，直接复制的结果到编程软件出现大量报错，于是将报错内容结合具体场景描述重新传输给 GAI，如“这段代码在 Python3.9 环境下运行出现语法错误，提示‘indentation error’，帮我修正并说明原因”，这种修正并非简单的复制粘贴，而是通过精准描述错误特征、限定运行环境，引导工具聚焦具体问题。更重要的是，使用者在“报错 - 反馈 - 修正”的过程中，被迫理解错误成因，将工具输出的半成品转化为学习素材，客观上训练了先拆解问题、再精准描述情况、最后调整应对策略的解决能力。这种整合不仅提升了交互效率，更在反复实践中内化为“遇事多问一句‘真的吗’、动手多查一遍‘对不对’、发现问题及时‘改一改’”的思维习惯，使批判性思维从分散的技巧提升为系统的认知能力。

4.3. 元认知能力的迭代

元认知能力的迭代本质上是大学生在与 GAI 的交互中，不断深化对“自身认知过程”“工具特性”“策略有效性”的反思与调控，具体表现为认知负荷协调、认知策略反思、AI 特性批判三个维度。这种自我调节在批判性思维塑造中起到关键作用，使批判性思维摆脱对具体任务的依赖，升华为可迁移的认知监控能力，最终实现从“被动适应”到“主动建构”的思维成熟。

在认知负荷协调方面，大学生对自身信息处理能力的动态感知与交互节奏调节，本质上是基于认知负荷理论对认知资源的主动调控。该理论指出认知负荷可分为内在负荷，即任务本身复杂度，外在负荷，即信息呈现方式引发的负荷，与相关负荷，即促进学习的有效负荷，三者需保持平衡以避免认知超载。GAI 若一次性输出冗长、未结构化的领域概述，如混杂概念、无关案例等，会迫使学生耗费资源在信息筛选、逻辑重组上。A21 曾因同时“要求 GAI 完成论文选题、框架、文献综述”导致接收的信息庞杂无序，内在负荷与外在负荷叠加，超出工作记忆即时处理容量，此时认知资源被大量消耗在低效信息整理上，相关负荷被严重压缩。反思后，建立在大学生对自身一次性能处理的信息复杂度、注意力持续时长的清醒认知上，通过预先定义信息框架，强制 GAI 输出符合认知预期的层次化信息，减少对次要内容的加工负荷。通过预判任务难度与自身负荷的匹配度，提升交互效率，避免思维资源的无效消耗，这种元认知策略直接反哺提示设计的精准性。

在认知策略反思层面，大学生逐步从“无意识使用”转向“有意识优化”交互策略，通过复盘与 GAI 的互动过程，主动评估自身指令设计、信息筛选、结果验证等环节的有效性。A35 最初直接让 GAI 生成完整文献综述，却因自身对核心文献缺乏基础认知，在课堂提问中难以解释综述逻辑；A27 在实验心理学课程中完全依赖 GAI 归结实验方法，直到汇报时被老师指出方法错误才意识到“自己没学好相关知识，根本看不出 GAI 在胡编乱造”。这些经历促使他们反思策略本质，GAI 的有效性并非绝对，其输出质量既受限于训练数据的覆盖范围与算法逻辑的固有局限，也依赖于使用者自身的知识储备与校验能力。由此，大学生的策略调整从简单调整操作步骤，转向深入思考如何通过与 GAI 的协同，更好地实现自身认知目标，这种转变背后是对“工具辅助”与“自身认知”关系的重新定位。通过持续对比不同策略的效果，大学生逐渐形成适配自身思维特点的交互模式，将批判性思维从应对具体问题转为优化认知过程的自觉意识。

在 AI 特性认知与批判层面，大学生逐渐突破对技术黑箱的认知，形成对 GAI 功能边界、可靠性局限与依赖风险的理性判断。从功能适配性的精准把握来看，大学生已能根据任务特性匹配工具能力。A14 根据任务特性选择工具，“写代码，我更喜欢用通义，如果是写文章或者润色文字，我就会选文心一言”，

A47 通过交互实践“对不同 GAI 工具的表现有更深入的认识”，进而形成针对性的使用策略，当涉及特定任务时，会直接调用在该领域表现更优的工具，体现了对工具功能边界的理性判断；大学生对 GAI 的认知并未止步于功能优势的识别，而是进一步延伸至对其可靠性局限与使用风险的批判性审视。这既包括对输出内容真实性的警惕，也涵盖对过度依赖可能导致的自身能力退化的预判。A21 在提交学校作业时，会主动避免过度依赖 GAI，以此防范学术抄袭等潜在风险，A10 认为“GAI 对于人们来说只起到一个拐杖的作用”，人应当把握关键的决策问题。这种认知不仅揭示了 GAI 输出的可靠性局限，更触及人与智能工具协同的边界，标志着对 AI 的认知从“技术崇拜”转向“理性驾驭”，将对工具的批判性认知内化为自身认知体系的一部分。

4.4. 负面交互案例

尽管前述研究发现揭示了大学生在与 GAI 交互中通过提示设计、过程调节和元认知迭代提升批判性思维的积极路径，但研究也观察到并非所有交互都能导向积极结果。部分学生未能有效利用提示工程，其批判性思维未得到显著提升，甚至出现了对工具的过度依赖或能力退化的风险。

部分学生未能实现从模糊指令到系统性提示设计的跨越。他们持续使用宽泛、模糊的初始提问，即使得到泛化或低质量输出，也缺乏动力或能力进行指令的迭代优化。这种停滞状态表现为一种追求快速获取答案的倾向，学生将提示工程视为额外负担，缺乏深入思考和精心设计指令的耐心，例如，受访者 A3 使用 GAI 时“问了几次感觉它答得都不太对，我就懒得再细问了，要么换别的方式查，要么随便用点能用的。”其深层原因主要在于两方面：一是元认知反思不足，未能有效审视自身指令的缺陷，或无法将 GAI 的低质输出归因于提问的模糊性，倾向于简单归咎工具“不好用”或“能力有限”；二是提示设计知识与技能欠缺，不了解如何构建有效的提示，或缺乏将复杂问题分解、结构化的必要能力。

部分学生将 GAI 视为思考外包的工具，直接依赖其生成完整答案、论证甚至决策建议，自身不再进行深入的解释、分析、推论和评估。受访者 A22 “在小组讨论中复述 GAI 观点作为己见，但当被进一步追问时却无法回答”。他们依赖于 GAI 快速生成看似合理答案的能力，将任务目标从“学习与理解”异化为“完成任务/获取答案”，忽视了批判性思维训练本身的价值。而当学生自身在特定领域知识储备或基本分析能力严重不足时，他们更容易全盘接受 GAI 的输出，缺乏识别错误或进行校验的知识基础，导致“越不懂，越依赖；越依赖，越不懂”的恶性循环。

这些无效或负面的交互案例表明，提示工程提升批判性思维的有效性受到学生个体的认知投入度、元认知能力、信息素养、学习动机以及对 GAI 工具批判性认知水平等多重因素的显著制约。部分学生因认知惰性、技能缺失、过度信任或目标错位，未能激活或反而抑制了自身在解释、分析、推论、评估和自我调节方面的认知过程，甚至产生了对工具的依赖，存在能力退化的风险。这些发现强烈提示，单纯依靠学生与 GAI 的自然交互并不足以确保批判性思维的发展，系统性的教育干预至关重要，以帮助学生克服惰性、识别风险、掌握方法，从而真正将 GAI 转化为批判性思维训练的“脚手架”而非“替代品”。

5. 教学活动实施建议

为将前述研究的理论发现转化为可操作的教学路径，针对性破解不同学科中批判性思维培养的差异化需求，以下结合文学评论与实验报告撰写两个典型学科场景，设计具体教学活动案例，为大学生借助 GAI 交互提升批判性思维提供实践指引。

5.1. 文学评论

文学评论注重文本解读的深度、情感共鸣的真实性及论证逻辑的严谨性，与 GAI 交互时，需引导学生通过精准提示设计、深度过程调节及元认知反思，强化文本分析、多维论证等批判性思维能力。活动

主题以经典文学作品的现代性解读“《呐喊》中阿 Q 形象的当代映射分析”为例，活动流程分为四个阶段：

第一阶段为初始提示生成。让学生自主设计初始提示，如“分析《阿 Q 正传》中阿 Q 的形象”，并基于 GAI 的输出，分组讨论其局限性，重点关注是否存在解读泛化、语言“人机味”过重而缺乏文学性或未结合原著细节等问题，参照研究中“论证质量核验”逻辑，明确需优化的方向。

第二阶段为提示优化。引导学生结合“多角色论证”策略，将提示修正为“请你作为一名文学评论家，聚焦阿 Q 的‘精神胜利法’，从‘个体心理与群体行为的关联’角度，分析这一形象在当代网络舆论场中的映射；需引用原著中具体情节，比如挨打后骂娘、画圈时的自欺这类内容”。此过程需让学生明确，设定角色能让解读更具专业性，明确分析的维度可让分析方向更聚焦，要求结合原著细节作为论据，能增强论证的实证性。

第三阶段为过程调节与人工验证。基于优化后的提示获取 GAI 输出后，组织学生进行验证：一方面核验逻辑性，检查“精神胜利法”的当代映射分析是否与原著中阿 Q 的行为逻辑一致；另一方面进行实证可靠性验证，通过核对原著，确认引用情节的准确性。若发现 GAI 输出存在与原著矛盾等问题，参照“结果动态修正”策略，进一步细化提示，如补充“需具体说明‘精神胜利法’在当代的 3 种表现形式，比如网络争论中的自欺式辩解这类情况，并对应原著情节分析其共通性”。

第四阶段为元认知复盘。引导学生反思两次提示设计的差异，分析优化后的提示如何通过明确角色、划定分析维度、提出论据要求，推动 GAI 输出更具深度；同时反观自身在设计提示时，是否预先完成了思维结构化，是否规避了仅求结果而忽视过程的思维惰性。

评估标准围绕三个维度展开：一是提示词的精准度，即是否明确角色定位、分析维度及论据要求，能否有效引导 GAI 输出聚焦且深入的内容；二是分析的深度与文学性，即基于 GAI 输出形成的评论，是否实现了文本细节与当代映射的有机关联，情感表达是否贴合文学作品的人文内涵，论证逻辑是否严谨；三是过程调节的有效性，即学生能否通过评估 GAI 输出的局限，主动修正提示并验证结果。

通过这一活动，学生可逐步掌握文学评论场景下与 GAI 交互的批判性策略，既避免对工具输出的盲从，又能以提示设计为抓手，强化自身对文本的深度解读与论证建构能力，实现批判性思维在文学领域的针对性提升。

5.2. 实验报告撰写

实验报告的核心在于操作过程与结果分析的紧密关联，其批判性思维的培养不仅依赖逻辑推导，更需扎根于实验操作的规范性、现象观察的细致性及结果与操作的对应性。与 GAI 交互时，需引导学生将工具辅助与实际实验操作深度结合，通过操作细节记录、现象差异分析、结果归因验证的连贯过程，避免脱离实操空谈理论或依赖 GAI 生成标准化结果的风险，强化基于实验过程的推理、操作规范性的反思及结果差异的实证分析能力。活动主题以经典生物学实验“叶绿体色素的提取与分离”为例，活动流程贯穿实操记录、问题诊断、交互修正三个维度，分为四个阶段：

第一阶段为初始提示尝试与预实验问题暴露。让学生自主设计初始提示，如“写一份叶绿体色素的提取与分离实验报告”，基于 GAI 的输出，分组开展预实验，即简化步骤，仅完成提取与初步分离。通过预实验发现 GAI 方案的实操缺陷：例如 GAI 建议研磨叶片时加无水乙醇，却未说明用量，实际加 5 mL 导致提取液过稀；或层析液倒入烧杯至液面没过滤液细线，导致色素溶解在层析液中，未出现分离带。结合研究中过程调节的逻辑，明确提示需补充操作细节量化，如试剂用量，以及关键步骤禁忌，如层析液液面高度等要求。

第二阶段为提示优化与实验方案细化。引导学生结合预实验结果优化提示：“作为实验操作者，基于预实验中研磨时无水乙醇过量导致提取液浓度低、层析液没过滤液细线导致分离失败的问题，设计完

整实验方案。需参考《基础生命科学实验指导》中色素提取的试剂配比规范，并解释这些参数设置的原因”。此过程需让学生明确，GAI 的方案需经实际操作检验，提示设计必须融入实验中用量、顺序、禁忌等实操细节，报告撰写的依据是具体操作而非抽象理论。

第三阶段为实操结果核验与动态修正。学生按优化后方案完成完整实验，获取真实结果，例如观察到滤纸条上仅出现 3 条色素带，缺胡萝卜素带，或叶绿素 a 带颜色偏黄，可能因研磨时未加碳酸钙导致叶绿素被破坏。随后将实验现象与操作记录输入 GAI，要求“请基于本次实验结果，如缺胡萝卜素带、叶绿素 a 带色浅，分析可能的原因，需结合实验操作过程，如研磨时间、试剂添加情况”。若 GAI 输出仅列举理论原因，而忽视实操细节，则引导学生参照人工验证机制修正提示，需结合本次实验的具体操作；并说明补做实验的设计思路，即延长研磨时间至 5 分钟，重新提取分离，以验证操作对结果的影响。此环节突出理科以操作过程解释结果的特点，GAI 的理论分析再全面，也需服从实际操作记录，批判性思维体现在对操作细节与结果差异的主动关联，而非盲从工具的标准化结论。

第四阶段为元认知复盘与实验反思。引导学生对比两次提示与实验的关联，分析优化后的提示如何因融入预实验的操作失误而更具实操性；反思 GAI 的理论分析与实际结果的偏差，探讨是否因提示未明确研磨时间量化要求导致 GAI 忽视操作时长对结果的影响；最终明确实验报告的撰写核心是用操作过程解释结果，用结果反推操作规范性，与 GAI 交互的关键是让工具聚焦于实验事实的分析，而非生成脱离实操的理想结果。

评估标准围绕三个维度展开：一是提示词的关联性，即是否包含实验材料用量、关键步骤禁忌、操作时长等细节，能否引导 GAI 输出贴合实际操作的分析框架；二是报告的实证对应性，即基于 GAI 辅助形成的报告，是否将色素带颜色、宽度等观察结果与研磨时间、试剂添加等操作细节一一对应，分析原因时是否兼顾理论解释与实操失误；三是过程调节的针对性，即学生能否通过实际实验发现 GAI 方案的操作漏洞，并用具体操作记录修正工具的理论化输出。

通过这一过程，可帮助学生掌握结合实验操作细节设计提示、依托真实观察结果校验输出、关联操作失误分析结果差异的交互策略，通过 GAI 辅助与实际实验的协同，提升实验报告中操作描述的精准性、现象分析的实证性，以及对操作规范与结果关联性的研判能力。

6. 结语

本研究立足提示工程视角，通过对 50 位不同专业、年级本科生的半结构化访谈，系统探究了 GAI 交互对大学生批判性思维能力的影响机制。研究发现，提示设计的动态演进、过程调节的深度赋能及元认知能力的迭代构成关键影响路径，揭示了提示工程如何将 GAI 从潜在的思维弱化工具转化为批判性思维发展的有效载体。理论上，这一发现丰富了 GAI 与批判性思维关系的研究维度，拓展了提示工程在教育场景的理论内涵；实践中为高校利用 GAI 培养批判性思维提供了具体路径，即引导学生掌握提示设计策略、强化过程调节意识、提升元认知反思能力。未来研究可进一步结合纵向追踪实验，考察提示工程训练对批判性思维长期发展的影响，并探索不同学科背景下提示策略的适配性差异。同时，可关注提示工程与智能教育平台的融合应用，探索构建动态化、个性化的批判性思维训练体系，探究提示工程在更广泛学习场景中促进思维能力提升的普适性规律，为智能化时代人类认知能力的持续发展提供理论支撑与实践范式。

基金项目

江苏省学位与研究生教育教学改革重点课题(JGKT24_B047)，江苏省研究生科研与实践创新计划项目(KYCX25_3515)。

参考文献

- [1] 荀渊. ChatGPT/生成式人工智能与高等教育的价值和使命[J]. 华东师范大学学报(教育科学版), 2023, 41(7): 56-63.
- [2] Dwivedi, Y.K., Kshetri, N., Hughes, L., Slade, E.L., Jeyaraj, A., Kar, A.K., *et al.* (2023) Opinion Paper: “So What If ChatGPT Wrote It?” Multidisciplinary Perspectives on Opportunities, Challenges and Implications of Generative Conversational AI for Research, Practice and Policy. *International Journal of Information Management*, 71, Article ID: 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- [3] 黄峻, 林飞, 杨静, 等. 生成式 AI 的大模型提示工程: 方法、现状与展望[J]. 智能科学与技术学报, 2024, 6(2): 115-133.
- [4] 冯雨夙. ChatGPT 在教育领域的应用价值、潜在伦理风险与治理路径[J]. 思想理论教育, 2023(4): 26-32.
- [5] 李艳, 许洁, 贾程媛, 等. 大学生生成式人工智能应用现状与思考——基于浙江大学的调查[J]. 开放教育研究, 2024, 30(1): 89-98.
- [6] 孙文芳, 吴泳锟, 林承德, 等. 基于人工智能的散打项目动作识别与自动评分方法研究[J]. 体育学研究, 1-13. <https://doi.org/10.15877/j.cnki.nsic.20250605.001>, 2025-08-20.
- [7] 王春丽, 陈艳艳, 顾小清, 等. 协作学习中生成式人工智能促进创造性潜能的机理研究[J]. 电化教育研究, 2024, 45(11): 76-83.
- [8] 金云波, 龚盼盼, 包莹莹, 等. 强人工智能时代大学生自主学习能力发展面临的机遇与挑战——基于 ChatGPT 的审思[J]. 信阳师范学院学报(哲学社会科学版), 2024, 44(1): 105-111.
- [9] 夏亮亮, 沈可心, 王宇, 等. 生成式人工智能情境下学生学习能动性: 现状与影响因素[J]. 现代远程教育, 2025(2): 37-47.
- [10] 彼得·范西昂, 都建颖, 李琼. 批判性思维: 它是什么, 为何重要[J]. 工业和信息化教育, 2015(7): 10-27, 41.
- [11] 何军, 姚雯皓. 人工智能对大学生批判性思维培养的价值及策略研究[J]. 中国医学教育技术, 2024, 38(6): 746-750.
- [12] 韩晴晴, 陈春梅. 基于生成式人工智能的大学生批判性思维发展探究[J]. 深圳信息职业技术学院学报, 2025, 23(2): 1-8.
- [13] 孟天广, 李珍珍. 治理算法: 算法风险的伦理原则及其治理逻辑[J]. 学术论坛, 2022, 45(1): 9-20.
- [14] 方海光, 王显闯, 洪心, 等. 面向 AIGC 的教育提示工程学习提示单设计及应用[J]. 现代远程教育, 2024(2): 62-70.
- [15] 周由. 提示工程视角下 AIGC 如何赋能教学设计[J]. 中国信息技术教育, 2025(11): 93-98.