

DeepSeek使用强度与大学生学习适应 ——基于遮掩效应与归因方式的实证研究

王 彤, 李嘉睿, 席晓诗

天津商业大学理学院, 天津

收稿日期: 2026年8月17日; 录用日期: 2026年9月9日; 发布日期: 2026年9月18日

摘 要

在生成式人工智能深度渗透高等教育的背景下, 本研究聚焦DeepSeek使用强度与大学生学习适应之间的复杂关系, 引入AI自我效能感作为遮掩变量、归因方式作为个体差异变量, 构建“使用强度→自我效能感→学习适应”的遮掩效应模型。以514名使用DeepSeek的大学生为样本, 采用问卷调查与分层回归分析, 结果表明: (1) DeepSeek使用强度对学习适应具有显著负向直接效应, 验证了技术替代效应的存在; (2) AI自我效能感在二者之间起遮掩中介作用, 间接效应(正向)与直接效应(负向)符号相反, 控制中介变量后负向效应增强, 揭示自我效能感既缓冲技术替代冲击、又掩盖过度使用危害的双重属性; (3) 归因方式的组间比较显示, 内部归因者表现出显著更高的使用强度与自我效能感, 但学习适应水平与外部归因者无显著差异, 呈现“高使用-高自信-未发现显著优势”现象。研究拓展了技术接受模型与自我效能感理论在AI教育应用情境中的解释边界, 为高校差异化AI素养教育提供了实证依据与干预靶点。

关键词

DeepSeek使用强度, 学习适应, AI自我效能感, 遮掩效应, 归因方式

DeepSeek Usage Intensity and University Students' Learning Adaptation

—An Empirical Study Based on the Masking Effect and Attributional Styles Scenarios

Tong Wang, Jiarui Li, Xiaoshi Xi

School of Science, Tianjin University of Commerce, Tianjin

Received: August 17, 2026; accepted: September 9, 2026; published: September 18, 2026

Abstract

Against the backdrop of generative artificial intelligence's deep penetration into higher education, this study investigates the complex relationship between DeepSeek usage intensity and college students' learning adaptation, introducing AI self-efficacy as a suppressing variable and attribution style as an individual-difference variable. A suppression effect model of "usage intensity → self-efficacy → learning adaptation" was constructed and tested with a sample of 514 college students who use DeepSeek, employing questionnaire surveys and hierarchical regression analysis. The results indicate that: (1) DeepSeek usage intensity has a significant negative direct effect on learning adaptation, confirming the existence of a technology substitution effect; (2) AI self-efficacy serves as a significant suppressor in the relationship between usage intensity and learning adaptation, with the indirect effect (positive) opposite in sign to the direct effect (negative), and the negative direct effect strengthening after controlling for the mediator, revealing self-efficacy's dual nature of both buffering the negative impact of technology substitution and masking the hidden costs of overuse; (3) between-group comparisons by attribution style show that internal attributors exhibit significantly higher usage intensity and self-efficacy, yet demonstrate no advantage in learning adaptation compared to external attributors, presenting a "high usage-high confidence-no significance advantages found" phenomenon. This study extends the explanatory boundaries of the Technology Acceptance Model and self-efficacy theory within AI educational contexts, providing empirical evidence and intervention targets for differentiated AI literacy education in higher education institutions.

Keywords

DeepSeek Usage Intensity, Learning Adaptation, AI Self-Efficacy, Suppression Effect, Attribution Style

Copyright © 2026 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技术的快速发展,以 DeepSeek、ChatGPT 为代表的 AI 工具已深度渗透高等教育领域,成为大学生学习活动的重要辅助手段。教育部《教育信息化 2.0 行动计划》¹明确提出推动人工智能与教育教学的深度融合,但技术赋能的边界与风险尚未得到充分关注。现有研究多聚焦于 AI 工具对学习效率的积极促进作用,忽视了过度依赖可能带来的认知能力退化风险。技术替代理论指出,当学习者过度依赖外部智能系统时,其自主认知加工机会将相应减少,长期可能导致元认知监控能力与批判性思维技能的削弱。因此,探讨 AI 工具使用强度与学习适应之间的复杂关系,具有重要的现实紧迫性。

AI 自我效能感作为个体运用 AI 技术完成学习任务的信心判断,在技术使用与学习成效之间扮演着关键角色。社会认知理论认为,自我效能感是行为改变的近端预测因子,能够激发学习动机、调节学习投入。然而,现有研究对自我效能感的作用机制存在分歧:一方面,高自我效能感可能促进技术接受与深度使用;另一方面,基于有限经验形成的“虚假效能感”也可能掩盖技术依赖的真实风险。这种矛盾提示,AI 自我效能感可能并非简单的中介变量,而是具有遮掩效应的复杂机制——既缓冲技术替代的负

¹http://www.moe.gov.cn/srcsite/A16/s3342/201804/t20180425_334188.html

面冲击，也掩盖过度使用的潜在危害。

更为重要的是，个体对技术使用的认知评价存在显著差异。归因理论指出，内部归因者倾向于将行为结果归因于自身能力与努力，这种归因风格驱动其主动探索新技术工具，形成更高的使用频率与技术自信。然而，高频使用与高技术自信是否必然转化为更好的学习适应，尚需实证检验。这种“高使用 - 高自信 - 未发现显著优势”的潜在悖论，揭示了个体差异在技术使用向学习成效转化过程中的关键作用。

基于上述分析，本研究以技术接受模型、社会认知理论与归因理论为框架，构建“DeepSeek 使用强度→AI 自我效能感→学习适应”的遮掩效应模型，并考察归因方式的分组差异特征，系统检验 AI 自我效能感的中介机制与遮掩效应，以期揭示 AI 学习工具影响学习适应的内在机理与个体差异，为高校差异化 AI 素养教育提供实证依据与干预靶点。

2. 文献综述与假设提出

2.1. DeepSeek 使用强度的概念界定与研究价值

DeepSeek 使用强度是指大学生在学习活动中使用 DeepSeek 生成式 AI 工具的频率、时长及功能探索程度的综合指标。随着生成式人工智能技术的快速发展，以 DeepSeek 为代表的大语言模型已深度渗透高等教育领域，成为大学生完成课程作业、代码编写、文献检索等学习任务的重要辅助工具。不同于传统的信息检索工具，DeepSeek 具有强大的自然语言理解与生成能力，能够直接输出结构化的知识内容与解决方案，这种“即问即答”的交互模式深刻改变着大学生的学习行为模式。

2.2. AI 情境下大学生学习适应的内涵与维度

学习适应是指主体根据学习环境及学习的需要，努力调整自我以达到与学习环境平衡的行为过程[1]。作为大学生心理适应的核心组成部分，学习适应涵盖学习动机、学习能力、学习态度、教学模式、环境因素五个维度，其中学习动机、学习能力、学习态度属于个体主观因素，教学模式与环境因素强调外在情境对学习的制约作用。

在生成式 AI 深度渗透的学习情境下，学习适应的内涵发生了新的延展：学习者不仅需要适应传统的课程要求、教学节奏与评价标准，还需调适与 AI 工具的协作关系——既要学会有效利用 AI 提升学习效率，又要避免过度依赖导致的自主学习能力退化。本研究采用冯廷勇等编制的大学生学习适应量表[1]，从学习动机(学习目标的明确性与学习主动性)、学习能力(知识获取与问题解决能力)、环境适应(对教学模式与学习资源的适应程度)三个核心维度测量学习适应水平。

2.3. AI 自我效能感的中介作用与理论基础

AI 自我效能感源于 Bandura (1977)提出的自我效能理论[2]，是指个体对自己运用人工智能技术完成特定学习任务的能力判断与信心程度。Bandura 将自我效能感定义为“个体对自己能否利用所拥有的技能去完成某项工作行为的自信程度”，强调其作为主体性因素对行为的驱动作用。随着人工智能技术的普及，技术自我效能感(Computer Self-Efficacy)逐渐成为自我效能感研究的重要分支，指个体对自己使用计算机技术完成特定任务的能力信念[3]。

2.4. 遮掩效应的机制特征与研究解释力

遮掩效应(Suppressing Effect)是中介效应的一种特殊形式，由温忠麟和叶宝娟(2014)系统界定[4]。当中介变量与自变量对因变量的影响方向相反，且控制中介变量后自变量的直接效应增强时，即构成遮掩效应。与传统中介效应(自变量→中介变量→因变量同向传递)不同，遮掩效应中中介变量“掩盖”了自变量对因变量的真实效应强度。

在本研究模型中, DeepSeek 使用强度对学习适应具有负向直接效应(技术替代效应), 但通过 AI 自我效能感产生正向间接效应(使用越多→效能感越高→适应越好)。由于两条路径方向相反, AI 自我效能感既缓冲了技术替代的负面冲击, 也掩盖了过度使用的真实损害, 形成典型的遮掩效应。识别这一机制对于准确评估 AI 工具的教育风险、设计针对性干预策略具有重要的方法论意义。

2.5. 归因方式的调节作用与异质性逻辑

归因方式是指个体对行为结果原因解释的倾向性差异, 由 Rotter (1966)提出的控制点理论发展而来 [5], 经 Weiner (1979)系统化为归因理论[6]。内部归因者倾向于将结果归因于自身能力、努力等可控因素; 外部归因者则倾向于归因于运气、任务难度等不可控因素。这种归因风格的差异深刻影响个体的动机、情绪与行为模式。

在技术使用情境中, 归因方式可能影响工具探索行为与心理资源的关系。内部归因者因将技术掌握视为能力提升的契机, 更可能主动探索 AI 工具功能, 形成更高的使用频率与技术自信; 但这种由归因驱动的积极探索未必自动转化为更好的学习适应, 可能因过度自信而忽视批判性审视, 或因高频使用导致认知负荷过载。相比之下, 外部归因者虽使用频率较低, 却可能通过其他补偿机制(如社会支持、求助行为)实现同等适应水平, 形成“低使用-同适应”的补偿模式。这种个体差异为理解技术使用效果的异质性提供了重要视角。

2.6. 研究假设

- H1: DeepSeek 使用强度对大学生学习适应具有显著负向直接效应。
- H2: AI 自我效能感在 DeepSeek 使用强度与学习适应之间起显著遮掩中介作用, 即间接效应与直接效应符号相反, 控制中介后负向效应增强。
- H3 (探索性假设): 内部归因者与外部归因者在 DeepSeek 使用强度、AI 自我效能感上存在显著差异, 但学习适应水平无显著差异。

3. 研究设计

3.1. 研究对象

采用方便抽样, 选取天津市大学生为研究对象。发放问卷 595 份, 回收有效问卷 580 份, 有效率 97.5%。剔除 66 名未使用 AI 工具者, 最终 514 名 DeepSeek 使用者纳入分析。使用者具体信息如表 1 所示。

Table 1. Basic information statistical table for the survey sample
表 1. 调查样本基本信息统计表

名称	选项	频数	百分比(%)
性别	女	383	74.5
	男	131	25.5
院校类型	非双一流高校	346	67.3
	双一流高校	168	32.7

样本性别与院校分布符合高校人文社科类专业特征, 具有一定代表性。

3.2. 研究工具

研究工具包括三个量表及归因方式单题项测量, 均采用 Likert 量表计分:

(1) DeepSeek 使用强度量表。自编量表, 含使用频率题项 1 题(采用 5 点计分)与功能探索度题项 5 题(采用 0/1 计分), 共计 6 个计分项, 合成总分作为使用强度指标。本研究中该合成指标的 Cronbach's $\alpha = 0.82$ (基于标准化后的题项得分计算)。

(2) AI 自我效能感量表。本研究采用 2 个题项测量大学生使用 DeepSeek 时的 AI 效能感。分别为: “使用 DeepSeek 后, 您对自己完成学习任务的信心变化是?” (1 = 显著降低, 5 = 显著提升)与“您使用 DeepSeek 学习新知识、完成课程作业的意愿如何?” (1 = 非常低, 5 = 非常高)。两题得分之和作为 AI 自我效能感的操作性指标(得分越高代表使用 AI 时的效能感越强)。本研究中该合成指标的内部一致性系数为 0.85 (两题相关系数 $r \approx 0.74$)。

(3) 学习适应量表。采用冯廷勇等编制的大学生学习适应量表[1], 根据本研究情境选取其中学习动机、学习能力、环境适应三个维度, 共 6 题项(1 = 完全不符合, 5 = 完全符合)。本研究中 Cronbach's $\alpha = 0.88$ 。

(4) 归因方式测量。本研究采用单题项对学生在 DeepSeek 辅助学习情境下的归因倾向进行探索性测量。题目为: “当 DeepSeek 帮助您解决复杂学习问题时, 您对自身学习能力的评价倾向于”, 选项 1 “认为自己能力得到验证” 定义为内部归因(编码 1), 选项 2 “部分归功于 AI” 与选项 3 “主要依赖 AI” 合并为外部/其他归因(编码 0)。本研究因条件限制采用单题项测量进行探索性二分化分组比较, 未来研究建议使用经过信效度检验的成熟量表例如结合学业归因风格问卷, 如 ASQ 进行情境化修订, 并将归因方式作为连续变量纳入调节效应或调节的中介效应分析, 以开展更精确的调节效应分析。本研究中内部归因者 109 人(21.2%), 外部/其他归因者 405 人(78.8%), 用于探索性分组比较。

3.3. 模型构建

基于技术接受模型、社会认知理论与归因理论, 构建“DeepSeek 使用强度→AI 自我效能感→学习适应”遮掩效应模型(见图 1)。假设 DeepSeek 使用强度对学习适应具有负向直接效应, 同时通过 AI 自我效能感产生正向间接效应, 二者方向相反形成遮掩。归因方式作为探索性分组变量, 用于考察内部与外部归因者在各变量上的差异特征, 不参与遮掩效应模型检验。鉴于当前单题项测量的局限性, 本研究仅能做初步的组间比较, 未能将其作为连续变量纳入模型检验调节效应, 这是本研究的一个重要局限。

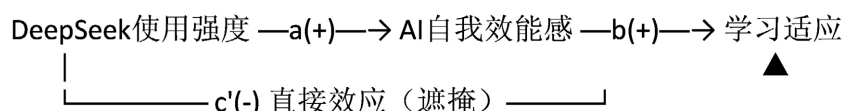


Figure 1. Masking effect model diagram

图 1. 遮掩效应模型图

3.4. 分析方法

采用 SPSS 26.0 进行描述性统计、独立样本 t 检验、Pearson 相关分析与分层回归。以归因方式为分组变量进行探索性 t 检验, 比较两组在各变量上的差异。运用 Hayes (2013)开发的 Process 插件(Model 4)检验遮掩效应[7], Bootstrap 抽样 5000 次, 置信区间 95%, 间接效应置信区间不包含 0 则效应显著。

4. 研究结果

4.1. 核心变量的总体分布特征

各核心变量的描述性统计结果如表 2 所示。样本 DeepSeek 使用强度均值为 5.06 (SD = 1.15), 处于中等使用水平; AI 自我效能感均值为 7.98 (SD = 1.13), 学习适应均值为 7.67 (SD = 1.61), 均高于对应量表 midpoint, 表明样本整体的 AI 自我效能感与学习适应水平处于中等偏上。各变量偏度绝对值均小于 2, 峰度

绝对值均小于 1，近似服从正态分布，满足后续参数检验的前提假设。

Table 2. Core variables-descriptive data (N = 514)
表 2. 核心变量描述性数据(N = 514)

变量	样本量 N	最小值	最大值	平均值	标准差	偏度	峰度
DeepSeek 使用强度	514	2.00	9.00	5.06	1.15	0.36	0.44
AI 自我效能感	514	4.00	10.00	7.98	1.13	0.44	-0.30
学习适应	514	2.00	10.00	7.67	1.61	-0.87	0.68
归因方式(0 = 外部/其他, 1 = 内部)	514	0.00	1.00	0.21	0.40	1.46	0.12

4.2. 不同归因方式下的组间差异特征

以归因方式(0 = 外部/其他, 1 = 内部)为分组变量，对核心变量进行独立样本 t 检验，结果如表 3 所示。

Table 3. Independent samples t-test results for grouped attribution methods
表 3. 归因方式分组的独立样本 t 检验结果

变量	莱文方差齐性检验		平均值等同性 t 检验			效应量	
	F	p	t	df	p (双侧)	Cohen d	95% CI
DeepSeek 使用强度	0.815	0.367	4.678	512	<0.001	0.51	[0.29, 0.73]
AI 自我效能感	40.693	<0.001	10.317	307.985	<0.001	0.81	[0.59, 1.03]
学习适应	3.615	0.058	1.205	512	0.229	0.13	[-0.08, 0.35]

注：AI 自我效能感因方差不齐，采用校正 t 检验结果；Cohen d ≥ 0.8 为大效应，0.5~0.8 为中等效应，<0.2 为极小效应。

研究发现，内部归因者在 DeepSeek 使用强度上显著更高(d = 0.51，中等效应)，AI 自我效能感亦显著更强(d = 0.81，大效应)，但两组学习适应水平无显著差异(d = 0.13，极小效应) [8]。因此，更为谨慎的表述是：内部归因者表现出“高使用、高自信却未发现显著优势”的特征。这一结果并不能直接证明内部归因者不存在学习优势，也不能据此断定认知负荷过载[9]或工具依赖是导致组间差异不显著的原因。统计功效不足、组间样本量不平衡、测量误差以及其他补偿性学习策略，都可能影响非显著结果的解释。未来研究应结合效应量、置信区间和等效性检验进一步判断两组在学习适应上的差异是否具有实际意义，并直接测量认知负荷、AI 依赖、输出验证行为和元认知监控等潜在机制。

4.3. 变量间的相关关系特征

对核心变量进行 Pearson 相关分析，结果如表 4 所示。

Table 4. Core variable correlation matrix
表 4. 核心变量相关矩阵

变量	1	2	3	4
DeepSeek 使用强度	1			
AI 自我效能感	0.202	1		
学习适应	-0.143	0.313	1	
归因方式(0 = 外部/其他, 1 = 内部)	+0.202	+0.336	+0.056	1

注：*p < 0.01 (双侧)，归因方式与学习适应的相关系数未达到显著水平(p = 0.229)。

DeepSeek 使用强度与 AI 自我效能感呈弱正相关($r = 0.202$), 但与学习适应呈弱负相关($r = -0.143$); AI 自我效能感与学习适应呈中等正相关($r = 0.313$), 三者相关模式符合遮掩效应的前提条件。

4.4. AI 自我效能感的遮掩效应检验结果

采用分层回归分析检验 AI 自我效能感在 DeepSeek 使用强度与学习适应间的遮掩效应, 结果如表 5 所示。

Table 5. Hierarchical regression results for the masking effect test
表 5. 遮掩效应检验的分层回归结果

变量	模型 1 (总效应 c)	模型 2 (a 路径)	模型 3 (直接效应 c')
	B(SE) [β]	B(SE) [β]	B(SE) [β]
(常量)	9.490 (0.466)	6.968 (0.327)	6.014 (0.484)
DeepSeek 使用强度	-0.190** (0.061) [-0.136]	0.203*** (0.043) [0.203]	-0.292*** (0.060) [-0.209]
AI 自我效能感	—	—	0.499*** (0.060) [0.350]
性别	-0.356 (0.161) [-0.097]	-0.038 (0.113) [-0.015]	-0.337 (0.151) [-0.091]
院校	-0.142 (0.149) [-0.041]	0.070 (0.105) [0.029]	-0.177 (0.140) [-0.052]
R ²	0.032	0.048	0.143
F 值	5.652***	8.563***	21.26***

注: * $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$; 括号内为标准误, []内为标准化系数 Beta; 模型 1 为使用强度对学习适应的总效应, 模型 2 为使用强度对 AI 自我效能感的预测, 模型 3 为加入 AI 自我效能感后使用强度对学习适应的直接效应。

H1 得到支持: DeepSeek 使用强度对学习适应具有显著负向预测作用($B = -0.190$, $p = 0.002$), 验证技术替代效应的存在。模型解释了有限但显著的变异(3.2%), 提示技术使用行为仅是影响学习适应的众多因素之一。

H2 得到支持: 遮掩效应的三项判定条件均成立。其一, a 路径显著: DeepSeek 使用强度正向预测 AI 自我效能感($B = 0.203$, $p < 0.001$), 模型解释了 4.8%的变异, H2a 成立。其二, b 路径显著: AI 自我效能感正向预测学习适应($B = 0.499$, $p < 0.001$), H2b 成立。其三, c'路径增强: 控制 AI 自我效能感后, DeepSeek 使用强度的负向效应绝对值增大($|\beta|$ 从 0.136 增至 0.209), H2c 成立。

综上, AI 自我效能感在 DeepSeek 使用强度与学习适应间起显著遮掩效应。间接效应为 0.1015 (95% CI [0.054, 0.159]), 与直接效应($B = -0.292$)符号相反, 揭示了一个关键悖论: AI 自我效能感作为心理资源, 既缓冲了技术替代的负面冲击, 也掩盖了过度使用的真实损害。

5. 讨论

5.1. 核心发现：技术替代效应与遮掩机制

本研究有三项核心发现。
(1) 技术替代效应的显现。DeepSeek 使用强度对学习适应的负向预测作用($B = -0.190$, $p = 0.002$)挑战了“技术赋能学习”的线性假设, 揭示生成式 AI 工具在学习场景中的潜在风险。根据认知负荷理论, 过度使用 DeepSeek 导致外在认知负荷超载, 学习者工作记忆资源被用于评估 AI 输出而非深度加工, 长期可能削弱元认知监控与批判性思维。但拓展了机制解释: 他们仅报告表层相关, 本研究揭示了遮掩效应的存在——AI 自我效能感掩盖了真实损害。

(2) 遮掩效应的识别。AI 自我效能感并非简单的中介变量,而是具有“双刃剑”属性的遮掩变量。其正向间接效应($a \times b = 0.102$)与 DeepSeek 使用的负向直接效应($c' = -0.292$)方向相反,形成“使用越多→效能越高→适应越好”的虚假传导路径。这一机制解释了为何部分高使用者自我感觉良好却实际适应不佳:自我效能感的心理补偿作用掩盖了认知能力退化的长期风险。

(3) 个体差异的探索性差异。分组比较发现,内部归因者的 DeepSeek 使用强度和 AI 自我效能感均高于外部/其他归因者,但在学习适应方面未发现显著优势。该结果只能说明,在本研究样本和测量条件下,更高的使用强度与更高的 AI 自我效能感尚未表现为显著更高的学习适应水平,不能直接据此推断存在“高使用-未能转化为高学习适应”的确定性机制。这一现象可能受到学习任务类型、使用质量、输出验证行为、认知负荷和元认知监控等因素影响,也可能与样本结构和统计功效有关。因此,后续研究需要通过多题项归因量表、客观学习表现指标以及行为实验,对上述可能机制进行直接检验。

5.2. 理论贡献:多重理论模型的边界拓展与机制补全

本研究的理论贡献在于三重拓展。

(1) 拓展技术接受模型的解释边界。传统 TAM 模型假设“感知有用性→使用意愿→使用行为→积极结果”的线性路径,本研究揭示遮掩效应的存在,证明同一心理变量(AI 自我效能感)可同时作为遮掩变量与中介变量存在,为模型修正提供实证依据。

(2) 丰富自我效能感理论的情境解释。Bandura 理论强调自我效能感的正向驱动作用[3],本研究补充其“暗面”——基于有限经验形成的“虚假效能感”可能掩盖行为风险,为理解“高自信-低成效”悖论提供机制解释。

(3) 提供个体差异的探索性视角。本研究将归因方式引入生成式 AI 辅助学习情境,发现内部归因者在使用强度和 AI 自我效能感方面具有较高水平,但本研究未发现其学习适应具有显著优势。这一结果为进一步考察归因方式与 AI 使用效果之间的异质性关系提供了初步线索,但由于归因方式采用单题项测量且仅进行了分组比较,尚不足以证明内部归因者存在“过度自信陷阱”,也不能将归因方式解释为已经得到验证的边界条件。

5.3. 实践启示:AI 辅助学习困境下的分层干预路径

基于上述发现,提出三点实践建议。

(1) 重塑 AI 素养教育理念。摒弃“使用越多越好”的认知误区,建立“有效使用+自主验证”的双轨模式,将批判性审视 AI 输出纳入核心素养。

(2) 差异化干预设计。对内部归因者,强化元认知监控训练,破解“能力幻觉”;对外部归因者,降低技术接入门槛,同时引导积极归因,避免“自我设限”。

(3) 建立使用质量评估机制。超越使用时长等表层指标,关注认知加工深度与自主验证频率,建立 AI 辅助学习的边界意识。

5.4. 研究局限与展望:研究设计约束与后续深化方向

本研究存在四方面局限。

(1) 横截面设计难以确立因果时序,未来可采用追踪研究或实验操纵验证遮掩效应方向。

(2) 样本来源和专业范围具有一定局限(局限于大数据专业),研究结果的外部效度仍需通过不同地区、院校类型、专业类别和年级的样本进行检验。

(3) 归因方式采用单题项测量,分类较为粗略,且样本中内部归因者占比偏低(21.2%),未来研究建议使用经过信效度检验的成熟量表,如改编版学术归因风格问卷 ASQ 进行测量,将归因方式作为连续变

量纳入模型,开展更精细的调节效应分析,而非简单的二分化分组比较。

(4) 本研究对认知负荷、AI 工具依赖、输出验证行为和元认知监控等可能机制未进行直接测量。因此,文中关于“高使用-未能转化为高学习适应”的解释仍属于理论推测,不能视为已经得到实证验证。未来研究可采用认知负荷量表、AI 依赖量表、学习过程记录、输出核查任务和行为实验,直接检验这些变量在 AI 使用与学习适应之间的作用。

(5) 本研究对 AI 自我效能感的测量采用了 2 题简版指标,主要基于 Q7(信心判断)与 Q10(使用意愿)合成,虽在操作上贴近效能感的定义,但题项覆盖不够全面,未来研究应采用经过严格信效度检验的多题项 AI 自我效能感量表进行重复验证。

6. 结论

本研究以 514 名大数据专业大学生为样本,系统检验了“DeepSeek 使用强度→AI 自我效能感→学习适应”遮掩效应模型,得出以下结论:

(1) 技术替代效应存在: DeepSeek 使用强度对学习适应具有显著负向预测作用,验证“高频使用不等于高效学习”的现实悖论。

(2) 遮掩效应显著: AI 自我效能感作为遮掩变量,既缓冲技术替代的负面冲击,也掩盖过度使用的真实损害,形成方向相反的间接路径,但该路径是否代表“虚假效能感”或具体的心理补偿机制,仍需结合效能判断准确性和实际学习表现进行验证。

(3) 个体差异凸显: 内部归因者呈现“高使用-高自信-未发现显著优势”特征,提示归因风格驱动的工具探索行为未必自动转化为实际学习成效。

(4) 理论模型拓展: 突破技术接受模型的线性假设,揭示 AI 学习工具影响机制的复杂性,为进一步解释不同学生使用 AI 工具后学习结果的差异提供了初步分析框架。

本研究为高校 AI 素养教育提供了实证依据与干预靶点: 在享受技术便利的同时,守护自主学习的核心空间。

7. 建议

本研究突破了对 AI 工具使用的单一正面认知,揭示了生成式 AI 背景下技术依赖、心理资源与个体差异的深层交互机制,丰富了人工智能教育应用领域的心理学研究,也为高校 AI 时代的人才培养、学生 AI 素养教育提供了实证依据。

7.1. 高校应正视技术依赖风险,构建系统化 AI 素养教育体系

(1) 树立科学的 AI 教育理念,以理性视角认知技术替代效应。

摒弃“AI 工具使用越多越好”的认知误区,认识到过度依赖作为挑战性压力可能带来的负面效应,通过专题讲座、课程引导等形式,帮助学生理解技术替代的本质与合理使用的边界,避免过度依赖引发自主学习能力退化。

(2) 针对不同归因风格学生差异化设计干预方案。

对内部归因者,需强化元认知监控与批判性思维训练,建立“AI 辅助+自主验证”的双轨学习模式,避免技术自信演变为盲目依赖;对外部归因者,应降低 AI 工具使用门槛,通过同伴互助、操作示范等方式提升其技术接入信心,同时引导其建立积极的归因模式,将技术探索视为能力提升的契机而非威胁。

(3) 搭建心理与学业融合的服务平台。

整合心理健康中心与教学管理资源,开展 AI 使用行为专项监测与辅导,提供个性化学业指导;建立

学生 AI 使用强度预警机制, 重点关注“高使用 - 未发现显著适应优势”群体, 及时干预过度依赖问题, 培养学生合理使用 AI 工具的意识与能力, 在享受技术便利的同时守护自主学习的核心空间。

7.2. 深化课堂教学改革, 激发批判性思维与元认知能力

(1) 重构 AI 辅助学习的课程设计模式。

教师应在课程设计中明确 AI 工具的使用边界, 设置“无 AI 辅助”的深度学习环节, 强制保留学生的自主认知加工空间; 通过项目式学习、翻转课堂等模式, 引导学生在使用 AI 工具后进行反思与验证, 培养批判性审视 AI 输出的能力。

(2) 建立 AI 自我效能感的科学培育机制。

AI 自我效能感的提升不应仅源于操作熟练度, 更应建立在“有效使用 - 真实提升”的良性循环之上; 教师需引导学生建立准确的自我评估, 避免因技术自信而陷入“能力幻觉”, 将自我效能感与元认知监控能力相结合。

(3) 破解“高使用 - 未发现显著适应优势”群体的转化困境。

针对 DeepSeek 使用强度高但学习适应水平低的内部归因者, 重点开展学习策略指导与认知负荷管理训练, 帮助其从“工具依赖型”学习向“策略调节型”学习转变; 针对低使用低适应的外部归因者, 通过归因训练与技术培训双管齐下, 提升其技术探索意愿与使用质量。

基金项目

天津市大学生创新创业训练计划项目(编号: 202510069086); 天津市普通高等学校本科教学质量与教学改革研究计划重点项目(编号: A251006902); 2026 年天津商业大学本科教学改革项目(编号: TJCUJG2026111)。

参考文献

- [1] 冯廷勇, 苏缙, 胡兴旺, 等. 大学生学习适应量表的编制[J]. 心理学报, 2006(5): 762-769.
- [2] Bandura, A. (1977) Self-Efficacy: Toward a Unifying Theory of Behavioral Change. *Psychological Review*, **84**, 191-215. <https://doi.org/10.1037/0033-295x.84.2.191>
- [3] Compeau, D.R. and Higgins, C.A. (1995) Computer Self-Efficacy: Development of a Measure and Initial Test. *MIS Quarterly*, **19**, 189-211. <https://doi.org/10.2307/249688>
- [4] 温忠麟, 叶宝娟. 中介效应分析: 方法和模型发展[J]. 心理科学进展, 2014, 22(5): 731-745.
- [5] Rotter, J.B. (1966) Generalized Expectancies for Internal versus External Control of Reinforcement. *Psychological Monographs: General and Applied*, **80**, 1-28. <https://doi.org/10.1037/h0092976>
- [6] Weiner, B. (1979) A Theory of Motivation for Some Classroom Experiences. *Journal of Educational Psychology*, **71**, 3-25. <https://doi.org/10.1037/0022-0663.71.1.3>
- [7] Hayes, A.F. (2013) Introduction to Mediation, Moderation, and Conditional Process Analysis: A Regression-Based Approach. Guilford Press.
- [8] Cohen, J. (1988) Statistical Power Analysis for the Behavioral Sciences. 2nd Edition, Erlbaum.
- [9] Sweller, J. (1988) Cognitive Load during Problem Solving: Effects on Learning. *Cognitive Science*, **12**, 257-285. https://doi.org/10.1207/s15516709cog1202_4