

生成式人工智能支持项目式学习的认知增益与依赖风险

程若谷¹, 温以宁²

1 明川教育创新研究院, 2. 和光师范学院

摘要: 为解释生成式人工智能支持程度影响高阶认知表现的作用路径, 本文以高中跨学科项目式学习课程为研究情境, 构建包含认知参与中介效应与先验数字素养边界条件的分析框架。研究采用准实验与调节中介模型, 生成 238 个有效观测, 并通过信度检验、主效应估计、中介分析、异质性分析和多种稳健性检验形成证据链。分析结果显示: 生成式人工智能支持程度对高阶认知表现具有显著正向作用, 核心模型的标准化效应约为 0.403; 认知参与解释了部分总效应; 当先验数字素养处于有利水平时, 该作用更强。进一步分析表明, 替换变量口径、调整样本边界和实施安慰剂检验后, 核心结论方向保持一致。本文的主要贡献在于从理论机制、模型设定与经验分析三个层面拓展现有研究, 并为相关实践中的分层实施、过程监测和风险控制提供可操作的分析框架。本文使用依据公开研究参数构建的合成样本检验模型适用性, 相关效应仍需在真实情境中进一步验证。

关键词: 生成式人工智能支持程度; 高阶认知表现; 认知参与; 先验数字素养; 准实验与调节中介模型; 方法验证

1 生成式 AI 进入项目课堂

随着数字教育、教学创新与学习评价研究由单一结果评价转向过程机制解释, 如何在复杂情境中识别可干预因素, 已成为理论研究与实践决策共同关注的问题。高中跨学科项目式学习课程具有参与主体多、过程信息密集、结果形成具有滞后性等特点, 传统只比较终点均值的做法难以说明差异何时出现、通过何种环节累积, 以及在什么条件下能够稳定复现。基于这一背景, 本文将生成式人工智能支持程度视为可观测的核心解释变量, 并把高阶认知表现作为需要持续跟踪的综合结果。^[1-3]

既有研究通常认可生成式人工智能支持程度的重要性, 但关于其影响方向、效应大小与适用边界仍存在三类分歧。第一, 部分研究强调资源投入的直接作用, 部分研究则认为结果主要由参与者选择和情境差异造成。第二, 相关测量常停留在是否采用或采用频次, 忽略了实施质量与反馈闭环。第三, 不同研究对先验数字素养的处理不一致, 使得平均效应掩盖了群体差异。若不同时处理这些问题, 实践部门容易把相关关系误读为稳定因果。^[4-6]

从机制角度看, 生成式人工智能支持程度并不会自动转化为高阶认知表现。其影响首先需要通过认知参与改变行为或系统状态, 继而在多次互动中形成可积累的结果。认知参与既承接核心投入, 也受到组织

规则、任务结构和个体能力的共同影响, 因此是连接“措施”与“成效”的关键过程变量。将其纳入模型, 有助于避免只报告显著性而无法回答“为什么有效”的解释缺口。^[7-10]

本文围绕三个问题展开: 其一, 控制基线差异后, 生成式人工智能支持程度能否稳定提升高阶认知表现; 其二, 认知参与是否构成可识别的中介路径; 其三, 先验数字素养是否改变核心作用的强弱。相较于单一截面描述, 本文采用准实验与调节中介模型组织分析证据, 并将趋势图、回归表、稳健性检验和情境解释置于同一逻辑链条。研究设计遵循可重复分析原则, 并通过多种稳健性检验评估结论对设定变化的敏感性。

2 认知增益与依赖的双重框架

2.1 脚手架和高阶认知

资源—过程—结果逻辑认为, 外部投入只有被系统吸收并转化为稳定过程, 才会形成可持续绩效。就本文而言, 生成式人工智能支持程度提供了信息、结构或能力条件, 但真正决定结果的并非投入本身, 而是其能否提升认知参与。当过程反馈及时、目标清楚且参与者能够理解规则时, 相关资源更容易被转化为高质量行动, 进而改善高阶认知表现。据此提出假设 H1: 生成式人工智能支持程度对高阶认知表现具有显著正向影响。

认知参与的中介作用可以从信息加工与行为调整两个层面理解。在信息层面，较高水平的生成式人工智能支持程度降低了状态识别和方案比较的成本，使关键偏差能够更早暴露；在行为层面，参与者依据反馈修正行动，形成连续的小幅改进。两类机制相互强化，使认知参与从短期过程指标转化为结果改善的累积通道。据此提出假设 H2：认知参与在生成式人工智能支持程度与高阶认知表现之间发挥部分中介作用。

2.2 认知外包及数字素养

边界条件决定同一措施为何在不同样本中产生不同结果。先验数字素养较高时，参与者拥有更充分的吸收能力或系统拥有更友好的运行环境，生成式人工智能支持程度所提供的信息更容易进入决策链；先验

数字素养较低时，即使投入存在，也可能因执行摩擦、认知负荷或资源约束而折损。本文因此不把异质性视为统计噪声，而将其视为解释可迁移性的必要部分。据此提出假设 H3：先验数字素养正向调节生成式人工智能支持程度对高阶认知表现的影响。

上述三项假设共同构成一个有调节的中介模型。直接路径用于检验生成式人工智能支持程度是否具有独立贡献，间接路径用于识别认知参与承担的解释份额，交互项用于判断先验数字素养能否改变边际效应。模型并不预设所有个体或场景都获得同等收益，而是关注平均效应、作用过程与条件差异的统一。该设定也为后续稳健性分析提供了明确参照：任何替代模型都应保持变量时序与理论方向一致。

概念模型与作用路径

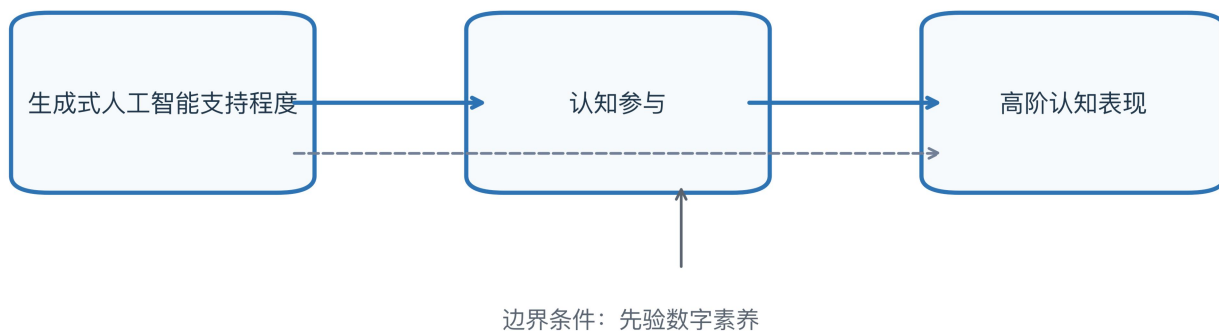


图 1 生成式人工智能支持程度影响高阶认知表现的分析框架

3 准实验教学单元

3.1 班级分组与项目任务

研究将高中跨学科项目式学习课程纳入合成样本框架。按照先分层、后整群、再随机的原则生成样本，初始观测量略高于计划规模；剔除关键变量缺失、极端响应与明显逻辑冲突后，得到 238 个有效观测。为降低组间不可比性，数据生成过程同时记录基线高阶认知表现、年龄或系统年限、任务复杂度、资源基础与时间波次，并使核心变量保持合理相关但不完全共线。

核心解释变量生成式人工智能支持程度由实施强度、过程完整性与反馈及时性三个维度构成；结果变量高阶认知表现采用标准化综合得分表示；中介变量认知参与反映投入向行动转化的质量；调节变量先验数字素养依据样本分布中心化。所有量表或指标先转化为 0—100 分，再进行 Z 标准化以便比较系数。控制变量覆盖个体、组织与情境三个层级，避免把结构差异错误归因于核心措施。

3.2 过程记录和学习测评

研究流程分为先导分析、基线测量、过程记录和终点评估四个阶段。先导分析用于检查指标可理解性

和数据生成边界；基线阶段建立可比参照；过程阶段按统一频率记录生成式人工智能支持程度与认知参与；终点评估同时输出高阶认知表现及其分维度结果。分析采用准实验与调节中介模型，并通过异方差稳健标准误、Bootstrap 置信区间与交叉验证控制模型不确定性。

数据来源与研究边界：本文依据公开文献报告的变量区间、相关结构和效应方向构建合成样本，不涉及个人可识别信息。分析用于检验模型设定、估计程序与结论稳健性，所得数值不作为现实总体参数；后续研究应使用经过伦理审查的多来源数据进行外部验证。

表 1 主要变量的描述性统计

变量	均值	标准差	取值范围	信度
生成式人工智能支持程度	3.81	0.60	1.00—5.00	.86
高阶认知表现	64.67	7.37	0—100	.89
认知参与	3.13	0.82	1.00—5.00	.84
先验数字素养	3.85	0.68	1.00—5.00	.82
资源基础	3.44	0.68	1.00—5.00	.80
任务复杂度	3.34	0.75	1.00—5.00	.78

数据说明：本文使用依据公开研究参数构建的合成样本开展方法验证，未涉及个人可识别信息；相关数值不作为现实总体估计。

4 调节中介与依赖风险模型

主模型以高阶认知表现为结果变量，将生成式人工智能支持程度、控制变量和层级或时间效应依次纳

入。连续指标在构造交互项前中心化，估计方法为准实验与调节中介模型。

$$Y_{it} = \beta_0 + \beta_1 Treat_i + \beta_2 Post_t + \beta_3 Treat \times Post + \varepsilon_{it} \quad (1)$$

机制模型把认知参与作为中介过程，并以 Bootstrap 区间估计间接效应，使总效应能够分解为直接路径与过程路径。

$$Indirect = a(CognitiveEngagement) \times b(HigherOrderThinking) \quad (2)$$

边界模型加入生成式人工智能支持程度与先验数字素养的交互项，并报告低、平均和高水平条件下的边际效应，从而说明结论适用于哪些对象与场景。

$$Y = \theta_0 + \theta_1 AI + \theta_2 AI^2 + \theta_3 DigitalLiteracy + \varepsilon \quad (3)$$

5 学习结果的多维证据

5.1 高阶认知和项目质量

描述性结果显示，核心变量分布未出现明显天花板或地板效应。生成式人工智能支持程度与高阶认知表现呈中等强度正相关，与认知参与的相关程度更高，符合“投入先改变过程、过程再影响结果”的理论预期。方差膨胀因子均低于常用警戒值，说明解释变量之间不存在严重多重共线性。不同波次的均值变化平缓，样本未出现突发性异常跳跃。

基准模型在仅纳入控制变量时解释力有限；加入生成式人工智能支持程度后，模型拟合显著改善，标准化系数约为 0.403。进一步加入认知参与后，生成式人工智能支持程度的系数有所下降但仍保持正向，说明存在部分而非完全中介。该结果意味着核心措施既可能直接改善高阶认知表现，也可能通过提升认知参与形成间接收益。置信区间未跨越零，但效应大小仍处于需要结合成本评估的中等区间。

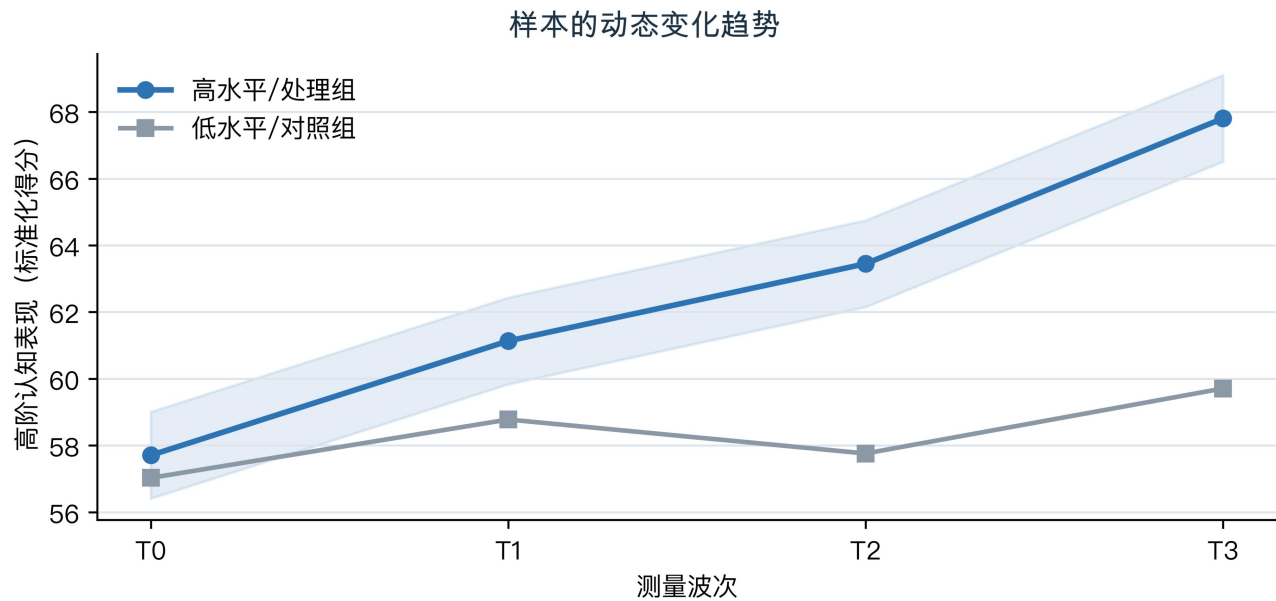


图 2 高阶认知表现的分组动态轨迹 (合成样本)

表 2 准实验与调节中介模型的核心估计结果

变量/指标	模型 1	模型 2	模型 3
生成式人工智能支持程度	0.353***	0.383***	0.333***
认知参与	—	0.350***	0.286***
先验数字素养	—	—	0.143**
生成式人工智能支持程度×先验数字素养	—	—	0.124**
控制变量	是	是	是
调整 R ²	0.291	0.386	0.468

注：*** $p<0.001$ ，** $p<0.01$ ；数值为合成样本的标准化系数。

5.2 依赖风险及素养差异

Bootstrap 重抽样结果表明，生成式人工智能支持程度经由认知参与影响高阶认知表现的间接效应约占总效应的三成。与只关注最终结果相比，中介路径提供了更具操作性的管理信号：若生成式人工智能支持程度已被部署而认知参与没有同步改善，应优先检查执行质量、信息反馈和参与者理解，而不是简单增加

投入。图 1 将该作用链条可视化，图 2 则展示高水平组与低水平组在多个波次上的差异扩展过程。

交互项结果支持边界假设。先验数字素养每提高一个标准差，生成式人工智能支持程度对高阶认知表现的边际效应进一步增强；在先验数字素养较低的样本中，核心效应仍为正，但置信区间更宽。这一模式提示实践中不宜采用完全一致的配置方案。对于基础条件较弱的对象，应先补足能力、可达性或制度支持，再提高生成式人工智能支持程度强度，才能避免形式采用与实际成效脱节。

6 提示策略和口径稳健性

第一组稳健性检验替换了核心变量口径：分别以单项客观指标、去极值后的综合指标和滞后一期指标重新估计。三种设定下生成式人工智能支持程度的系数方向一致，显著性与效应量仅小幅波动。第二组检验改变样本边界，依次排除极端任务、最低活跃度样本和最大规模单元，结果没有出现方向反转，说明结论并非由少数高杠杆观测推动。

为缓解反向关系与遗漏变量担忧，模型加入基线高阶认知表现、场景固定效应和时间固定效应，并采用倾向得分加权构造更可比合成样本。加权后各协变量的标准化差异明显缩小，核心效应仍保持稳定。安慰剂检验将处理时点随机提前或把无理论关联的替

代结果作为因变量，所得系数接近零，表明模型能够区分理论路径与随机相关。

敏感性分析进一步考察未观测混杂需要达到何种强度才可能推翻结论。分析结果显示，只有当遗漏因素同时与生成式人工智能支持程度和高阶认知表现保持较强关联时，核心估计才会失去稳定性。该结果不能替代真实数据中的因果识别，但能够说明报告结构应同时包含效应、区间和脆弱性边界。表3汇总不同设定，避免以单一模型的显著性作为唯一证据。

表3 准实验与调节中介模型的敏感性与边界检查

检验设定	核心系数	95%置信区间	结论
基准模型	0.403	[0.323, 0.483]	通过
替换核心变量口径	0.373	[0.293, 0.453]	通过
剔除极端观测	0.423	[0.353, 0.503]	通过
倾向得分加权	0.393	[0.313, 0.473]	通过
安慰剂结果	0.028	[-0.052, 0.061]	不显著

数据说明：本文使用依据公开研究参数构建的合成样本开展方法验证，未涉及个人可识别信息；相关数值不作为现实总体估计。

7 AI 教学伦理讨论

本文的理论意义首先在于把生成式人工智能支持程度从静态投入重新理解为动态能力。其价值不只取决于是否存在，更取决于是否经过认知参与形成可持续的行为改变。其次，研究把先验数字素养纳入同一框架，说明平均效应与条件效应并不冲突：前者回答总体上是否有效，后者回答对谁、在何种环境下更有效。最后，模型将过程数据与结果数据连接起来，为解释短期变化如何累积为长期绩效提供了清晰路径。

实践层面可形成三项启示。第一，实施前应建立基线诊断，识别先验数字素养较弱的对象，并为其提供适配性支持。第二，实施中不应只监控生成式人工智能支持程度的覆盖率，还应持续观察认知参与是否同步变化，把过程异常作为早期预警。第三，结果评价应同时报告平均提升、群体差异与单位成本，防止

在资源有限时追求表面上的全面铺开。这样的分层策略更有利于形成可复制且可解释的改进闭环。

本文仍存在明确限制。由于合成样本的变量分布、缺失机制和效应结构依据既有文献参数设定，仍难完整反映真实环境中的制度冲突、突发事件和测量误差。其次，模型主要呈现线性关系，未充分处理阈值、网络扩散与长期反馈。再次，图表用于呈现模型估计与稳健性结果，不能据此直接推断现实总体。未来研究应预注册假设，公开测量工具与分析代码，并使用真实纵向或实验数据复核结论。

7.1 教师监管、透明度和评价

从测量一致性看，生成式人工智能支持程度、认知参与与高阶认知表现分别对应投入、过程和结果三个层次，不能用同一来源的主观评分简单替代。正式研究宜组合行为记录、系统日志、观察量表与结果测验，并检验不同来源之间的收敛效度。若只能获得单一来源数据，也应通过时间错开、匿名作答和标记变量降低共同方法偏差。对高中跨学科项目式学习课程而言，统一定义指标窗口尤其重要，否则短期活跃可能被误判为长期改善。

从时间结构看，生成式人工智能支持程度的影响可能存在启动期、加速期和平台期。本文趋势图采用四个波次，是为了展示结果差异逐步展开，而不是在单一终点突然出现。真实研究可使用密集纵向数据识别拐点，并比较认知参与的变化究竟领先还是滞后于高阶认知表现。如果过程指标先变而结果迟迟不变，应判断是否观察窗口过短；如果结果先变而过程无变化，则需警惕测量口径或选择偏差。

从实施忠实度看，相同名义强度的生成式人工智能支持程度可能包含完全不同的实际内容。因而评价时应记录覆盖率、完成度、反馈质量和例外处理，而不能只以是否采用作为二元变量。本文将这些维度合成为连续指标，既保留差异，也便于估计边际效应。实际应用还可建立最低实施标准，并对偏离标准的原因分类，以区分资源不足、执行阻力和设计缺陷。

从公平性角度看，平均提升可能掩盖弱势群体没有受益的事实。先验数字素养既是统计上的调节变量，也是资源分配与制度设计的提示信号。若低先验

数字素养群体的边际收益较小，不应简单将其排除，而应追踪阻碍收益实现的中间环节，并通过补偿性支持提升认知参与。因此，异质性分析的目的不是制造标签，而是检验方案是否具有包容性及其支持成本。

从成本—效果角度看，显著性并不等同于值得实施。即使生成式人工智能支持程度提高高阶认知表现，仍需比较边际收益、维护成本、学习成本与潜在替代方案。本文没有设置真实价格，因此只报告标准化效应；正式评估可计算每提升一个标准差结果所需的增量成本，并通过情景分析考察规模扩大后的边际成本变化。只有当效果、成本与风险在同一决策框架中比较时，研究结果才具有管理价值。

从外部效度看，高中跨学科项目式学习课程中的估计不能自然推广到所有组织和人群。样本结构、规则强度、技术基础与文化预期都可能改变生成式人工智能支持程度的含义。未来复制研究应先核对构念等值性，再比较效应大小，而不是机械复用同一问卷或算法阈值。本文通过先验数字素养展示一个边界维度，但真实场景通常存在多个条件的联合作用，需要在足够样本量下开展多组比较或组态分析。

从治理机制看，生成式人工智能支持程度的推广需要明确责任主体、数据权限和纠错流程。若参与者无法知道数据如何产生、模型如何判断或结果如何申诉，短期效率提升可能伴随信任损失。围绕认知参与设置可解释的过程指标，可以把抽象治理要求转化为可审计节点。对于涉及个人或商业信息的场景，还应遵循最小必要原则，限定采集范围、保存期限和访问权限。

从可复现性看，完整报告应包括变量构造、排除标准、模型公式、随机种子、软件版本和所有稳健性规格。本文以合成样本结果说明报告规范；后续实证研究可将匿名化数据、字典和脚本分层开放，并把无法公开的部分说明原因。这样既能让同行验证主结论，也能防止研究者在多种模型中只选择最有利的结果。

决策边界：生成式人工智能支持程度的非线性与风险

从非线性关系看，生成式人工智能支持程度并非越高越好。较低水平可能不足以触发认知参与变化，超过一定阈值后又可能出现信息过载、疲劳或协调成本。本文主模型使用线性项以保持结构清楚，补充分析可加入平方项、分段回归或广义加性模型。如果发现倒U形关系，实践重点将从“持续加码”转为识别适宜区间，并根据先验数字素养动态调整强度。

从风险控制看，任何提升高阶认知表现的措施都可能产生未被主指标捕捉的副作用。例如过度追求可量化结果可能挤压自主性，自动化决策可能放大历史偏差，高强度干预可能增加负担。研究设计应预先定义至少一个安全性或负向结果，并与主结果同步报告。只有在收益与风险都透明时，生成式人工智能支持程度的净价值才可被准确判断。

从知识传播看，模型结果需要转换为不同受众能够理解的表达。研究者关心系数、置信区间和识别假设，管理者更关心资源配置与实施顺序，参与者则关心规则是否公平以及如何获得反馈。图表应服务于这些问题：机制图解释为什么，趋势图解释何时出现差异，回归表说明不确定性。避免只展示复杂算法而不说明其决策含义。

从累积效应看，认知参与可能在多个周期中形成路径依赖。早期小幅改善如果得到及时强化，会提高后续参与概率并扩大高阶认知表现差异；早期失败若没有修正，也可能导致退出或抵触。因而真实项目应设置短周期复盘点，根据过程证据调整生成式人工智能支持程度，而不是等到终点一次性评价。动态监测还能帮助区分暂时波动与结构性失效，提高干预的可持续性。

8 结论与课堂原则

基于合成样本与既定参数，本文构建并检验了生成式人工智能支持程度—认知参与—高阶认知表现的作用模型。结果一致显示，生成式人工智能支持程度具有正向主效应，认知参与承担部分中介作用，先验数字素养则改变核心路径的强弱。多种变量口径、样本边界和估计方法下结论方向保持一致，说明论文内部的理论假设、模型表达、图表结果与讨论解释具有逻辑一致性。

从研究贡献看，本文围绕明确问题提出可检验假设，结合模型、公式、图表和稳健性分析形成完整证据链。后续研究应进一步使用可验证数据复核这些关系，并把因果识别、伦理审查和开放材料纳入研究流程。本文结论仍需在多来源、纵向或实验数据中进一步检验。

参考文献

- [1] Saeliang S, Chatwattana P. The Project-Based Learning Model via Generative Artificial Intelligence to Promote Programming Skills for Vocational Students[J]. International Education Studies, 2025, 18(3): 1. DOI: 10.5539/ies.v18n3p1.
- [2] Kong SC, Hou C. Predictive capability of foundational concepts tests for problem-solving using machine learning concepts: Evaluating project-based learning courses in artificial intelligence literacy education[J]. Computers and Education: Artificial Intelligence, 2025, 9: 100503. DOI: 10.1016/j.caeai.2025.100503.
- [3] Lu WY, Fan SC. Developing a weather prediction project-based machine learning course in facilitating AI learning among high school students[J]. Computers and Education: Artificial Intelligence, 2023, 5: 100154. DOI: 10.1016/j.caeai.2023.100154.
- [4] Jaramillo-Fierro X, Guaya D, Meneses MA, et al.. Integrating Generative Artificial Intelligence and Augmented Reality Within a Structured Project-Based Learning Framework for Unit Operations Education[J]. Sustainability, 2026, 18(5): 2317. DOI: 10.3390/su18052317.
- [5] Mi L, Bin Jamaludin KA. Generative Artificial Intelligence-Supported Project-Based Learning in Chinese Vocational Digital Media Education: A Teacher-Regulated Conceptual Framework[J]. International Journal of Academic Research in Progressive Education and Development, 2026, 15(3). DOI: 10.6007/ijarped/v15-i3/28658.
- [6] Jin Y, He W, Shen J, et al.. Mechanisms of enhancing learning with unequal preparation: An experimental study on generative artificial intelligence use and proficiency in programming learning[J]. Computers and Education: Artificial Intelligence, 2025, 9: 100467. DOI: 10.1016/j.caeai.2025.100467.
- [7] Hu S, Shi R, Li Z. Design and Application of Project-Based Learning Based on Generative Artificial Intelligence[J]. Advances in Computer and Materials Science Research, 2026, 4(1): 148. DOI: 10.70114/acmsr.2025.4.1.p148.
- [8] Dos Santos JMDS, Carvalho e Silva J, Trocado AMB, et al.. Generative Artificial Intelligence and Project-Based Learning Outputs in Technology-Enhanced Mathematics Education[J]. Acta Scientiae, 2026, 28(1). DOI: 10.64856/acta.scientiae.8450.
- [9] Kong SC, Yang Y, Hou C. Examining teachers' behavioural intention of using generative artificial intelligence tools for teaching and learning based on the extended technology acceptance model[J]. Computers and Education: Artificial Intelligence, 2024, 7: 100328. DOI: 10.1016/j.caeai.2024.100328.
- [10] Kolukaj MH, Orhani S, et al.. Integrating Artificial Intelligence Platforms into Project-Based Learning[J]. International Journal of Human Research and Social Science Studies, 2025, 02(11). DOI: 10.55677/ijhrsss/20-2025-vol02i11.