

生成式 AI 赋能下高校学生学术行为意愿与内在机制研究

何思佳¹, 卜胤丹¹, 周燕^{2*}

1. 华南农业大学 国际教育学院, 广东 广州 510642

2. 华南农业大学 数学与信息学院, 广东 广州 510642

DOI:10.61369/ASDS.2025060015

摘 要 : 随着生成式 AI 技术在教育中的广泛应用, 大学生在学术任务中使用 AI 工具的行为及其内在机制成为研究热点。本文构建三阶段研究框架, 首先通过聚类分析识别使用行为差异, 其次基于“多维前因—动机中介—行为结果”路径构建结构方程模型, 揭示功能认知、社会影响、伦理意识等因素通过动机中介影响使用意愿。最后, 借助 LDA 主题建模对开放式反馈进行补充分析, 提取学生关注的风险与期望。结果表明大学生使用行为具有显著异质性, 使用动机起关键中介作用。建议高校开展分层引导, 加强伦理教育, 促进 AI 工具合理使用。

关 键 词 : 生成式 AI; 学术行为; 聚类分析; 结构方程模型; 文本挖掘; 动机机制

Research on College Students' Academic Behavior Intention and Underlying Mechanisms Empowered by Generative AI

He Sijia¹, Bu Yindan¹, Zhou Yan^{2*}

1.College of International Education, South China Agricultural University, Guangzhou, Guangdong 510642

2.College of International Education, South China Agricultural University, Guangzhou, Guangdong 510642

Abstract : With the growing integration of generative AI in educational contexts, there is an urgent need to systematically examine college students' usage behavior and its underlying mechanisms in academic tasks. This study adopts a three-stage research framework. First, cluster analysis is employed to identify heterogeneous patterns of AI tool usage. Second, a structural equation model is constructed based on a "multidimensional antecedents-motivational mediation-behavioral outcomes" pathway, revealing how factors such as functional perception, social influence, and ethical awareness affect usage intention through motivational mediation. Finally, LDA topic modeling is applied to open-ended responses as a supplementary analysis to extract students' primary concerns and expectations. Results demonstrate significant heterogeneity in usage patterns, with motivation playing a key mediating role. The findings suggest that universities should adopt differentiated guidance strategies, strengthen AI ethics education, and promote the rational use of generative AI tools.

Keywords : generative AI; academic behavior; cluster analysis; structural equation modeling; topic modeling; motivational mechanism

引言

自2022年 ChatGPT 发布以来, 生成式 AI 在全球教育领域快速普及, 深刻改变了知识获取、学习支持及科研方式(王思遥等, 2024)^[1]。截至2024年, 我国生成式 AI 用户达2.3亿, 高校应用增长突出, 美国58% 高校教师已将其融入教学(CNNIC, 李艳等, 2024)^[2]。同时, 技术应用多元化, 在知识获取、学习支持、科研辅助等方面都展现出显著赋能作用(刘凯, 2025)^[3], 但也带来“数据偏见”“学术幻觉”等伦理风险(UNESCO, 2023)。

然而, 生成式 AI 重塑大学生学习生态, 也引发对技术依赖与认知异化的担忧。一是学术任务中 AI 依赖明显, 20%-30% 的毕业论文中含 AI 生成内容(张仟煜等, 2024)^[4]; 二是批判性思维下降, 主动验证 AI 生成内容准确性者不足三成(李艳等, 2024)^[2]; 三是

基金项目:

广东省高等教育教学改革项目“基于创新能力培养的统计学专业数据分析实验课程改革探索与实践”(粤教高函[2024]9号516); 2023年大学生创新创业训练计划项目“生成式 AI 赋能下高校学生学术行为意愿与内在机制研究”。

作者简介:

何思佳, 华南农业大学国际教育学院, 本科生, 研究方向为数据挖掘与智能决策;

卜胤丹, 华南农业大学国际教育学院, 本科生, 研究方向为数据挖掘与智能决策。

通讯作者: 周燕, 华南农业大学数学与信息学院, 讲师, 研究方向为数据挖掘与智能决策。

学术伦理边界模糊, AI 检测规避行为频现(孙旭等, 2024)^[5]。面对 AI 技术带来的学习变革和伦理挑战, 亟需系统研究大学生使用行为、动机及心理机制。

当前关于生成式 AI 在高等教育领域中的应用研究, 已初步形成以技术接受、认知发展与伦理治理为核心的多维分析框架。技术接受模型(TAM)强调感知有用性与风险影响行为意愿(张池, 2023)^[6]。认知路径模型(SEM)揭示 AI 对创造力的双重作用(王思遥等, 2024)^[1]。伦理治理则提出“人机协同”框架, 聚焦使用规范、透明度及风险认知(UNESCO, 2023)。同时, 学者们关注学生群体差异与反馈: 李艳等(2024)^[2]通过混合方法揭示学科差异与使用习惯关联; 周子琦等(2025)^[7]指出 AI 依赖抑制高阶思维; Darvishi 等(2023)^[8]指出社会规范、同伴行为等外部因素显著影响使用动机。

尽管研究不断深化, 但仍存在样本片面、周期短与跨学科融合不足等问题。本研究在此基础上, 结合结构方程模型构建“多维前因—动机中介—行为表现”路径模型, 融合聚类与文本挖掘, 探讨大学生在 AI 使用的行为特征与心理机制。现有研究表明, 感知有用性、易用性及内容可信度与规范感知等因素显著影响学生的使用态度与动机(段锦云, 2025^[9]; Darvishi, 2023^[8])。因此结合生态系统视角, 从个体、组织与社会三层引入安全信任、社会影响等外部驱动因素, 并纳入学术伦理意识作为前因变量(Abbas, 2023)^[10]。同时, AI 焦虑理论认为, 对技术替代的担忧会抑制学生自主探索 AI 的内在动机(Wang et al., 2022)^[11]; 但工具性外在动机和自我效能感可缓解其影响, 为此模型引入内外动机作为中介路径, 用于承接前因与行为。基于自我调节学习理论, 具备较强目标管理与反思能力的学生更能合理使用 AI, 减少依赖并促进深度思维(王思远等, 2024)^[12], 故将实际使用行为表现作为结果变量。

一、研究设计

(一) 分析框架

本研究以大学生为研究对象, 聚焦其在学术任务中使用生成式 AI 工具的行为特征, 综合引入技术接受模型、计划行为理论与信息系统使用行为研究成果, 构建三阶段分析框架。首先, 基于功能认知、适应性、安全与信任度等关键变量设计量表, 采集问卷并通过 K-means 聚类构建用户画像。其次, 结合“多维前因—动机中介—行为结果”路径模型, 运用结构方程模型(SEM)检验前因变量通过使用动机的中介路径对使用行为倾向的直接与间接影响关系。最后, 结合文本挖掘分析开放问答数据, 补充个体态度与心理特征, 增强模型解释力与研究深度。研究旨在构建系统可解释的行为路径模型, 并提出教育引导策略。

(二) 研究对象

本研究的研究对象为全国范围内具有生成式 AI 工具使用经验或认知的高校大学生, 涵盖不同年级、地区与学科, 确保样本多样性与代表性。研究重点聚焦大学生在学术任务中使用生成式 AI 的行为特征, 通过行为倾向、情感态度与心理认知等维度挖掘其使用动因与群体差异。

(三) 数据来源

数据来源于基于技术接受模型(TAM)、计划行为理论(TPB)等理论设计的调查问卷, 结合大学生使用生成式 AI 工具的行为特征与心理机制设定变量。问卷分四部分: 背景信息、使用现状、变量量化测量(功能认知、适应性、安全与信任度、社会影响、伦理意识、使用动机与行为表现等)及开放式反馈, 样本采用分层三阶段抽样。第一阶段按东中西部及东北地区高校省份分布抽样; 第二阶段在选定省份内按高校类型(985、211、普通本科)分层抽样; 第三阶段在高校内根据年级与专业进行个体抽样, 部分采用整群抽样。

(四) 研究工具

问卷通过“问卷星”平台在线发放, 设有逻辑检测题筛除无效样本, 前期完成信效度检验, 确保量表内部一致性与测量合理性。数据分析主要采用 Stata、AMOS 与 Python 工具。首先使用 Stata 进行 K-means 聚类分析, 基于大学生学术任务中使用生成式 AI 工具的行为特征, 划分典型用户类型并构建差异化画像。随后, 利用 AMOS 构建结构方程模型(SEM), 检验“多维前因—动机中介—行为结果”路径中各变量的直接与间接效应。最后, 运用 Python 中的 jieba 与 gensim 对开放式反馈进行文本预处理、词频统计与 LDA 主题建模, 并结合 SnowNLP 开展情感分析, 提取学生态度特征与心理反应。

二、影响大学生使用生成式 AI 工具的行为动机

(一) 聚类目的与思路

为识别大学生在学术活动中生成式 AI 工具使用中的典型行为模式与态度差异, 本文选取对 AI 工具的了解程度、使用意愿、使用频率、依赖程度、效率认知等变量量化, 运用 K-means 算法进行聚类, 最终划分出四类用户群体。

(二) 聚类结果与用户画像

使用 Stata 进行 K-Means 聚类, 依据轮廓系数法确定最优簇数为 4 类, 按使用频率与依赖程度划分为深度依赖探索型、中度使用成长型、谨慎尝试观望型与基本不使用者。基于聚类结果, 绘制散点图进一步可视化分析。

图 1 显示, 不同群体在使用频率与依赖程度上区隔明显, 进一步验证了聚类的合理性。四类群体由右上到左下依次沿斜线分布, “深度依赖探索型”表现出高度活跃与深度依赖的使用特征, “中度使用成长型”具备稳定的使用习惯和中等依赖倾向; “谨慎观望型”使用水平较低但具备一定关注与尝试意愿, “基本不使用者”的使用意愿与行为均处于最低水平, 反映大学生在 AI 工具使

用行为上存在显著异质性，可据此开展更具针对性的行为机制分析与教育引导。为揭示各群体使用特征与偏好差异，进行可视化处理（图2），辅助理解行为分层与态度差异。

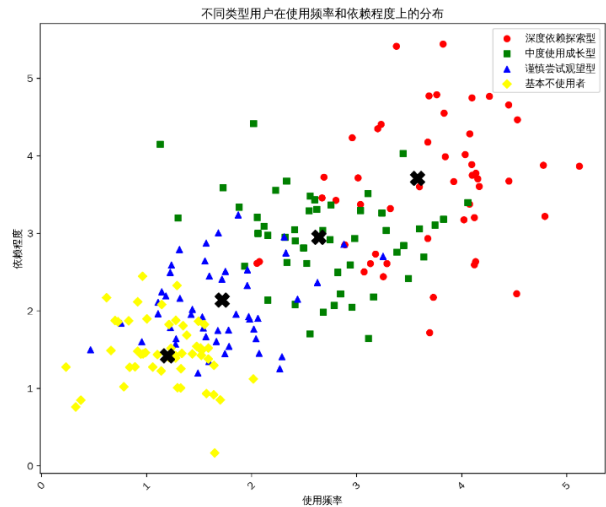


图1：不同类型用户在使用频率和依赖程度上的分布



图2：用户画像模型

用户画像显示，大学生在生成式 AI 的学术使用行为上呈现出明显的异质性。深度依赖探索型群体积极拥抱新技术，使用频率高；中度使用成长型使用稳定；谨慎观望型使用意愿有限，信任不足；基本不使用者接受度较低，普遍坚持传统学习方式。结果验证了聚类合理性，并为 AI 素养教育提供参考。

三、多维动因驱动下的大学生生成式 AI 使用行为机制建模与实证检验

（一）模型建立

1. 理论假设与模型构建

基于理论与前期调研，构建模型假设功能认知、适应性、安全与信任度、社会因素以及学术伦理意识均正向直接影响使用动机，且假设这五类因素通过使用动机中介对行为表现有间接正向

影响，使用动机显著正向影响行为表现。

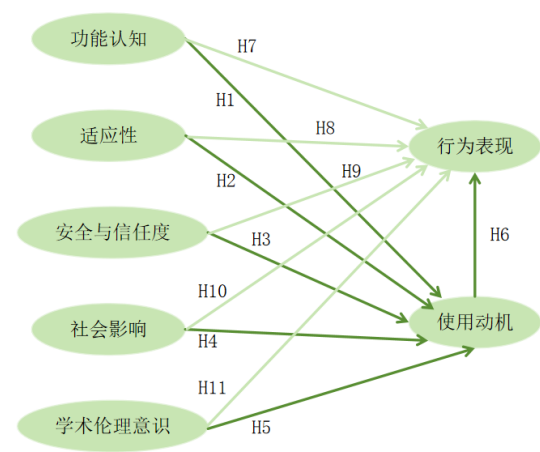


图3：结构方程模型假设

（二）数据处理

1. 测量维度设计

所有测量变量均采用 Likert 五级量表评估，具体见下表：

表1：Likert 五级量表

潜变量	测量变量
功能认知	学习效率感 PU1
	任务效率感 PU2
	成绩提升预期 PU3
适应性	使用便捷性 EA1
	理解难易度 EA2
	适应意愿 EA3
安全与信任度	内容信任度 TA1
	思维焦虑感 TA2
	能力退化担忧 TA3
社会影响	同伴影响力 SI1
	信息获取 SI2
	模仿倾向 SI3
学术伦理意识	违规规避意愿 AE1
	内容加工意愿 AE2
	诚信告知行为 AE3
使用动机	内在动机 IM1
	外在动机 IM2
行为表现	使用频率 BO1
	依赖程度 BO2
	持续使用意愿 BO3

2. 验证性因子分析

表2结果表明，问卷数据 KMO 值为 0.901，远高于 0.8，表明样本间共性强；Bartlett 球形度检验显著（ $p<0.001$ ），说明相关矩阵非单位矩阵。因此，样本适合进行因子分析。

表2：KMO 和巴特利特球形度检验

KMO		0.901
巴特利特球形度检验	近似卡方	1324.214
	自由度	136
	显著性	

为检验测量模型结构效度，进行验证性因子分析（CFA）。结果（表3）显示，CMIN/DF=2.486<5的可接受标准；RMSEA=0.059<0.08的判断标准，表明模型残差较小，拟合较好；GFI、IFI、TLI、CFI接近或达到0.90以上，分别为0.892、0.926、0.915和0.925，接近或达到理想标准，模型整体拟合较好。PNFI与PCFI分别为0.712和0.734，超过0.5的最低要求，说明模型简约且适配性强。

表3：验证性因子分析模型拟合指标

适配度指标	理想标准值	可接受标准值	研究结果	判别
CMIN/DF	≤0.05	≤5	2.486	合格
RMSEA	≤0.05	≤0.08	0.059	合格
GFI	≥0.90	≥0.70	0.892	合格
IFI	≥0.90	≥0.70	0.926	合格
TLI	≥0.90	≥0.70	0.915	合格
CFI	≥0.90	≥0.70	0.925	合格
PNFI	≥0.50	/	0.712	合格
PCFI	≥0.50	/	0.734	合格

3. 问卷信度与聚合效度分析

表4：各变量信度及聚合效度检验

变量名称	题项数	Cronbach's Alpha	AVE	CR
功能认知	3	0.821	0.515	0.852
适应性	3	0.713	0.563	0.811
安全与信任度	3	0.762	0.562	0.806
社会影响	3	0.770	0.531	0.788
学术伦理意识	3	0.810	0.551	0.801
使用动机	2	0.851	0.546	0.797
行为表现	3	0.834	0.522	0.788

为检验问卷信度与聚合效度，计算了Cronbach's α 、CR和AVE(见表4)。各变量的Cronbach's α 系数介于0.713至0.851，高于0.7的常规信度标准，表明各维度内部一致性良好，问卷整体信度较高。聚合效度方面，所有潜变量CR值在0.788至0.852之间(>0.7)，AVE值在0.515至0.563之间(>0.5)，均符合判断标准。上述结果表明各潜变量聚合效度良好，能有效反映各维度潜在结构特征。

4. 区分效度检验

表5：区分效度检验结果

	功能认知	适应性	安全与信任度	社会影响	学术伦理意识	使用动机	行为表现
功能认知	0.726						
适应性	0.492	0.744					
安全与信任度	0.524	0.516	0.761				
社会影响	0.411	0.418	0.432	0.739			
学术伦理意识	0.479	0.463	0.471	0.471	0.751		
使用动机	0.588	0.547	0.546	0.537	0.506	0.817	
行为表现	0.475	0.451	0.482	0.463	0.431	0.503	0.741

表5区分效度检验结果表明，四个潜变量的平方根AVE值（对角线）均大于其与其他潜变量之间的相关系数，符合For-

nell-Larcker 判别标准。说明各潜变量间区分度良好，量表在测量不同构念时具有良好辨别力。

（三）结构方程模型结果与路径分析

表6：结构方程模型路径系数分析

假设	路径关系	标准化系数	S.E.	P	结果
H1	使用动机 <--- 功能认知	0.483	0.079	0.001	接受
H2	使用动机 <--- 适应性	0.427	0.091	0.001	接受
H3	使用动机 <--- 安全与信任度	-0.253	0.081	0.001	接受
H4	使用动机 <--- 社会影响	0.301	0.065	0.001	接受
H5	使用动机 <--- 学术伦理意识	-0.201	0.053	0.001	接受
H6	行为表现 <--- 使用动机	0.774	0.105	0.001	接受
H7	行为表现 <--- 功能认知	0.201	0.072	0.061	拒绝
H8	行为表现 <--- 适应性	0.153	0.080	0.081	拒绝
H9	行为表现 <--- 安全与信任度	-0.115	0.052	0.093	拒绝
H10	行为表现 <--- 社会影响	0.098	0.073	0.155	拒绝
H11	行为表现 <--- 学术伦理意识	-0.059	0.042	0.232	拒绝

表6路径分析结果显示，H1至H6均达到统计显著水平。其中，“功能认知”（ $\beta=0.483$, $p<0.001$ ）、“适应性”（ $\beta=0.427$, $p<0.001$ ）、“社会影响”（ $\beta=0.301$, $p<0.001$ ）对“使用动机”均有显著正向影响，而“安全与信任度”（ $\beta=-0.253$, $p<0.001$ ）与“学术伦理意识”（ $\beta=-0.201$, $p<0.001$ ）则呈显著负向影响，表明功能优势与技术亲和力可增强动机，而安全顾虑与伦理意识则可能抑制其形成。此外，“使用动机”对“行为表现”具有显著正向作用（ $\beta=0.774$, $p<0.001$ ），验证其核心中介地位。相比之下，H7至H11中，前因变量对“行为表现”的直接路径均不显著（ $p>0.05$ ），其中，“功能认知”（ $p=0.061$ ）与“适应性”（ $p=0.081$ ）虽接近显著，但作用可能主要通过“使用动机”间接发挥。

表7：中介效应检验

中介路径	效应值	标准误	显著性	Bootstrapping	
				Bia-Corrected 95%CI	
				下限	上限
使用动机 <--- 功能认知	0.374	0.062	0.001	0.255	0.498
使用动机 <--- 适应性	0.330	0.068	0.001	0.210	0.463
使用动机 <--- 安全与信任度	-0.196	0.057	0.001	-0.310	-0.095
使用动机 <--- 社会影响	0.234	0.054	0.001	0.138	0.348
使用动机 <--- 学术伦理意识	-0.156	0.049	0.002	-0.258	-0.069

Bootstrapping 结果显示，功能认知、适应性与社会影响通过“使用动机”对行为表现产生显著正向中介效应（效应值分别为0.374、0.330和0.234，95%置信区间均不含0）；而安全与信任度、学术伦理意识的中介效应虽显著但为负（分别为-0.196和-0.156）。说明“使用动机”在各前因变量与行为表现之间起关键中介作用，既能强化正向影响，也可能传导负面认知对行为的抑制效应。

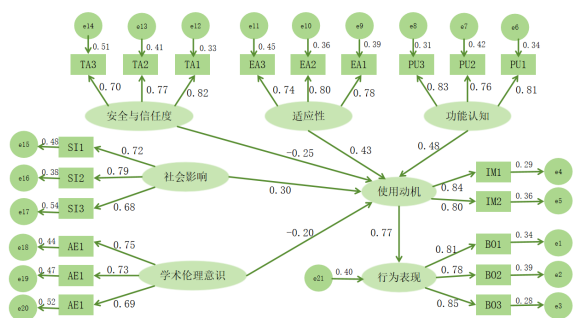


图4：结构方程模型结果

图4结构方程模型结果表明，整体模型较好验证了研究初设路径，虽部分路径未达预期，但拟合良好，各路径关系解释力强，能有效反映大学生学术任务中使用 AI 工具的动机结构与行为特征。

四、基于 LDA 模型的大学生学术活动中生成式 AI 使用主题挖掘

（一）数据的获取与处理

为深入探讨大学生在学术活动中使用生成式 AI 工具的行为特征、主观感受与潜在风险认知，本研究采用问卷与社交平台评论相结合的方式收集数据。一方面，收集531份结构化问卷，调查大学生在学术活动中生成式 AI 工具的使用情况；另一方面，通过 Python 爬虫从知乎、小红书、哔哩哔哩等平台获取相关评论，共收集有效评论11203条，去重清洗后保留10225条，用作 LDA 主题分析语料。数据经分句与分词处理，剔除冗余信息，转换为标准文本格式并进行分析。

（二）词频统计和词云图

表8：评论文本词频统计表

词语	频数	词语	频数
AI	3381	提高	1197
学术	2954	思路	1140
写作	2736	依赖	900
效率	1920	抄袭	855

高频词汇显示学生普遍关注 AI 在“写作”“效率”“提高”等方面的作用，认为是提升学术任务完成效率与质量的重要工具。同时，出现频率较高的“依赖”“抄袭”“独立思考”等有批判性的词反映出对 AI 可能带来的学术规范问题和思维惰性的担忧，揭示其对 AI 工具的双重态度。基于词频结果生成词云图，进一步可视化学生的关注重点。



图5：词云图

（三）评论文本的情感分析

本研究使用 Python 的 SnowNLP 工具对文本情感进行打分，得分区间为 [0,1]，高于0.5视为正向情绪，低于0.5为负向。情感倾向可反映用户的使用意愿，以直方图呈现结果。情感分析结果显示，大部分评论得分集中在0.8至1.0，反映出大学生普遍持积极态度，使用体验良好。由于情感分析难以揭示具体内容维度，后续引入 LDA 模型进一步挖掘主题与关注点。

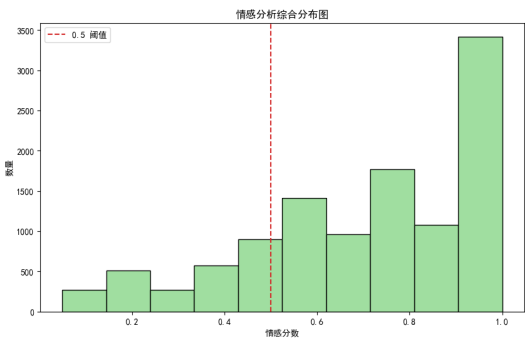


图6：情感分析综合分布图

（四）文本评论的 LDA 主题分析

本研究采用 LDA 主题模型对评论文本进行主题挖掘。LDA 基于“词袋”假设，将文本建模为多个潜在主题的组合，每个主题由若干高频词构成。在训练不同主题数的模型后，结合困惑度指标，最终确定最优主题数为4类，四个主题内容如下：

表9：大学生使用生成式 AI 主题词表

主题一	主题二	主题三	主题四
效率	抄袭	思考	提高效率
提高	造假	能力	思维
没有	容易	独立思考	论文
方便	侵权	思路	节省时间
更高	存在	高效	缺少
学习	重复	提升	依赖

由表9可知，主题一主要反映学生对生成式 AI 工具在学习过程持肯定态度，普遍认可其在提升学习效率与便利性的工具性价值，在任务完成与知识掌握方面展现积极作用。主题二反映了学生对于 AI 在学术活动中可能引发学术不端问题的关注，体现其对生成式 AI 使用边界和规范性的敏感与警觉。主题三关注生成式 AI 对学生思维能力的影 响，与学生对 AI 与自主学习关系的批判性理解相关，反映其在获得启发的同时，也反思依赖是否削弱学习的主动性与深度思考的能力。主题四体现生成式 AI 与学术任务完成的关联。体现 AI 在提升写作效率中的应用，同时揭示学生对其依赖与能力弱化之间的双重风险。

LDA 分析表明，大学生在重视 AI 效率的同时，也关注其对学术规范与思维能力的影 响，需加强伦理引导与能力培养，实现技术与成长的平衡，提升其在高校群体中的正向价值。

五、结束语

本研究结合聚类分析、结构方程模型与 LDA 主题建模，系统探讨了高校学生在学术活动中使用生成式 AI 工具的行为特征与

动因，得出以下结论：一是大学生在生成式 AI 使用上存在显著分层，专业背景与学习阶段是关键影响因素，其中理工科及高年级学生更倾向于高频使用。二是学生普遍持积极态度，认可 AI 工具的高效性与创新性，但对技术依赖、隐私泄露和学术伦理问题亦存明显担忧。三是使用动机在行为路径中起核心桥梁作用，功能认知显著提升动机水平，而信任度低、伦理风险感知等负向因素会抑制使用意愿。

基于此，提出三点建议：第一，高校应识别学生群体差异，精准推广 AI 工具。鼓励“高频用户”发挥示范作用，对“低接触群体”加强启蒙教育与实践引导。第二，强化 AI 伦理与风险教育，纳入通识课程，通过案例教学、规范制定等方式提升学生判断力与责任意识。第三，提升功能认知、减轻风险顾虑，借助课程培训和平台优化提升动机水平，并提供差异化支持，推动 AI 工具在学术场景中的可持续应用。

参考文献

[1] 王思遥, 黄亚婷. 促进或抑制：生成式人工智能对大学生创造力的影响 [J]. 中国高教研究, 2024, (11): 29-36.DOI:10.16298/j.cnki.1004-3667.2024.11.05.

[2] 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(01): 89-98.DOI:10.13966/j.cnki.kfjyyj.2024.01.010.

[3] 刘凯. 生成式 AI 工具在大学生学习模式创新中的应用研究 [J]. 大学教育, 2025, (03): 70-74+101.

[4] 张仟煜, 刘恺骁, 杨洁. AI 辅助或代写论文拷问大学的容忍边界 [N]. 中国青年报, 2024-07-08(005).DOI:10.38302/n.cnki.nzsqn.2024.002560.

[5] 孙旭, 钟秋菊, 张文涛. 生成式 AI 时代大学生智能学习助手：框架、挑战与应对 [J]. 终身教育研究, 2024, 35(04): 29-36+45.DOI:10.13425/j.cnki.jjou.2024.04.004.

[6] 张池. 大学生对于生成式人工智能工具的使用意愿研究——基于技术接受模型 [J]. 科技传播, 2023, 15(23): 131-135.DOI:10.16607/j.cnki.1674-6708.2023.23.034.

[7] 周子琦, 高飞, 方春晖, 等. 基于布鲁姆认知分类的大学生 AI 依赖风险分析与对策研究 [J/OL]. 云南民族大学学报 (自然科学版), 1-7[2025-04-10].http://kns.cnki.net/kcms/detail/53.1192.N.20250106.1122.002.html.

[8] Darvishi S. M., Hosseini S. H., & Mohammadi M. The Paradox of Student Agency in AI-Assisted Learning: An Experimental Analysis[J]. Journal of Educational Computing Research, 2023, 61(5): 1234 - 1256.

[9] 段锦云, 陈晓文, 李一飞. 基于 TAM 的高校教师生成式 AI 技术使用意愿研究 [J]. 现代教育技术, 2025(01): 102 - 108.

[10] Abbas J., Zhang W., & Mahmood S. Exploring the Academic Consequences of Generative AI Usage among College Students: A Mixed-Method Study[J]. Education and Information Technologies, 2023, 28(6): 7385 - 7404.

[11] Wang Y., Li H., & Xu Z. Understanding College Students' AI Learning Anxiety and Motivation: A Structural Equation Modeling Approach[J]. Computers & Education, 2022, 190: 104609.

[12] 王思远, 黄亚婷. 自我调节学习能力在生成式 AI 学习中的作用机制研究 [J]. 电化教育研究, 2024, 45(03): 81 - 89.