

应对生成式人工智能挑战：学术写作合规性框架与合法性边界界定

王书凡

安徽大学高等教育研究所，安徽 合肥

收稿日期：2025年5月22日；录用日期：2025年6月6日；发布日期：2025年7月31日

摘 要

随着AIGC技术的迅猛发展，其在学术写作中的应用日益广泛，既为个体进行学术写作提升带来了机遇，也引发了不当署名、学术剽窃，学术伪造和篡改等一系列学术不端的挑战，有必要探寻生成性人工智能在学术写作中运用的合法边界，对超出合法边界的行为进行认定和处理。本研究为规范AIGC在学术写作领域的合理应用、促进学术诚信体系建设提供了理论依据与实践指引。

关键词

AIGC，学术写作，学术不端，学术诚信，合法性边界

Addressing Challenges of Generative Artificial Intelligence: Compliance Framework and Legal Boundaries Demarcation in Academic Writing

Shufan Wang

Institute of Higher Education, Anhui University, Hefei Anhui

Received: May 22nd, 2025; accepted: Jun. 6th, 2025; published: Jul. 31st, 2025

Abstract

With the rapid development of AIGC technology, its application in academic writing has become increasingly widespread. This has not only brought opportunities for individuals to enhance their

academic writing capabilities but also triggered a series of challenges related to academic misconduct, such as improper authorship, academic plagiarism, fabrication, and falsification. It is necessary to explore the legal boundaries of using generative artificial intelligence in academic writing, and to identify and address behaviors that exceed these boundaries. This study provides theoretical foundations and practical guidance for regulating the rational application of AIGC in the field of academic writing and promoting the construction of an academic integrity system.

Keywords

Artificial Intelligence Generated Content, Academic Writing, Academic Misconduct, Academic Integrity, Legal Boundaries

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(Artificial Intelligence Generated Content, AIGC)技术的日渐兴起,我国政府于2023年7月发布《AIGC 服务管理暂行办法》(以下简称《暂行办法》),旨在把握技术变革机遇,促进AIGC的健康发展和规范应用[1]。AIGC以强大的知识生产能力为特征,基于算法和数据,能够较短时间内生成大量文本,甚至图片和短视频之类的内容,革命性的人机交互模式给社会生活很多领域带来了新的可能性,在教育领域给学术写作带来了颠覆性的影响。学术写作是在学术领域内,遵循特定规范与格式,运用专业语言,以客观、严谨、逻辑的方式阐述研究成果、论证观点并进行知识传播与学术交流的书面表达形式。一方面,个体在学术写作可以展现自己的知识,表达自己的立场,培养自己批判性思维的能力;另一方面,学术写作支持个体参与学习过程,学习构建知识的不同方式。AIGC通过模拟大脑的神经元,将人脑中的神经网络移植复制到人工智能中,即人工神经网络(ANN),使人工智能像人一样会思考和创作[2],这与个体的学术写作过程具有适切性。

以 ChatGPT 为代表的 AIGC 刚发布时,引起了政府,高校和行业组织的广泛抵制。顶级学术期刊 Science 主编即宣告其已成为一种文化轰动,且对科学界和学术界造成严重影响[3],并与 Nature 一同禁止将其列为论文作者。一些著名大学相继发布禁止学生使用 ChatGPT 完成作业、论文的规定,如英国包括牛津大学和剑桥大学在内的八所著名大学均规定使用 AIGC 完成作业论文是学术不端行为。学术界逐渐认识到 AIGC 在学术界的巨大潜力,在 AIGC 参与学术写作过程中保持谨慎探索的态度,开始逐渐探索 AIGC 在学术界的规范使用。国际出版社《科学》《自然》《细胞》《柳叶刀》《美国医学会杂志》等知名出版机构相继发布 AI 写作辅助指南,强调论文中必须标明 AIGC 辅助写作的内容且 AIGC 不具备作者资格。同时在我国,中国科技部发布的《负责任研究行为规范指引(2023)》[4],以及中国科学技术信息研究所公布的《学术出版中 AIGC 使用边界指南》,对于学术研究环节如何恰当使用 AIGC 做出了道德要求与学术规范,如禁止用 AIGC 生成项目申报书,AIGC 不能成为科研项目共同申报人或成果完成人等要求。2025年9月1日施行的《人工智能生成合成内容标识办法》(以下简称《标识办法》)要求服务提供者在生成的文本内容之起始或中间位置附加显式标识,辅助高校更为直观地追溯论文内容的真实来源,国内外的高校也相继发布了 AIGC 的使用辅助指南,探讨 AIGC 在高等教育中的合规使用。如一些美国高校要求学生在使用 AI 工具完成作业或论文时进行说明,并对使用 AI 的程度加以限制。我国华北电力

大学、湖北大学、福州大学、南京工业大学、天津科技大学等多所学校发布规定，明确学校将引入人工智能代写检测技术，抵制人工智能代写论文的行为。由此可以看出，AIGC 在学术写作的运用是具有合法性依据的。在我国，《人工智能生成合成内容标识办法》等规范性文件的出台，表明我国对于 AIGC 在学术写作的运用是抱有包容审慎的态度，对于 AIGC 在学术写作中的辅助性功能是鼓励和支持的。然而在实际的学术写作中，AIGC 在学术写作中的辅助功能模糊了学术不端的边界，容易产生不当署名，内容伪造和篡改以及学术剽窃的学术不端行为。我们应该如何在学术写作中正确规范使用 AIGC？AIGC 使用的合法使用边界在哪里？正是本文需要探讨的内容。

2. 文献综述

AIGC 在学术写作领域的影响是广泛且深远的，国内外很多学者都进行了探讨。从国内外已有的研究来看，学界认为 AIGC 对学术写作的影响主要是益处和可能的挑战两个方面。很多学者鼓励将 AIGC 应用到学术写作之中，主要考虑到以下方面。首先，AIGC 提高了学术写作的速度和质量。例如，中国科学院基于 ChatGPT 定制了一整套实用性功能，用于优化学术研究以及开发日常工作流程。该系统可以进行学术论文一键润色、语法错误查找、中英文快速互译、智能读取论文并生成摘要等工作。其次，在语言润色与多语种写作支持方面，针对非英语母语的学生，AIGC 能够提供实时的语法纠正、风格优化和术语标准化建议，显著提升学术写作的准确性和流畅性[5]。最后 AIGC 可以扩展个体学术写作的思路。通过根据学习者的写作样本实时生成反馈与建议，AIGC 能够实现针对性写作辅导[6]。例如，一些智能写作助手不仅可以指出文本中的语言问题，还能根据论文的逻辑性、论证深度等维度给出具体改进意见，从而促进学生的批判性思维与自我修正能力发展。

从可能带来的挑战来看，主要包括：第一，过度使用 AIGC 可能导致学生的创新能力和批判思维能力下降。第二，AIGC 生成的内容包含一些未经验证的错误信息，存在一定误导学生的可能。第三，AIGC 自动生成的论文本质上不会有真正的创新。此外 AIGC 对学术诚信的监管和学术评价也带来了挑战。一方面，传统的查重方式无法判断学位论文写作是否运用了 AIGC；另一方面，当前对 AIGC 参与论文写作的检测技术尚未成熟。有研究发现，32%的 ChatGPT 生成摘要被专家误识为原创，14%的原创摘要被误识为 ChatGPT 生成。值得注意的是，AIGC 工具的影响不可一概而论。如有研究指出 ChatGPT 的表现不同学科领域存在差异，在经济学领域表现优异，在计算机编程领域的表现也可圈可点，但在数学领域表现不佳[7]。另外，将 ChatGPT 应用于学术研究的不同任务或环节，效果也可能大相径庭。ChatGPT 在提供初步的研究想法、文献总结、撰写摘要等方面可能具有优势，但在文献引用、陈述研究问题、寻找研究缺口和数据分析等方面则差强人意，甚至会产生误导[8]。总体而言，AIGC 作为一种新兴智能技术，在学术写作中的辅助功能不容否认，但其合法性边界、伦理规范与应用风险亦需高度警惕与制度回应。

3. 生成式人工智能背景下合法性风险

随着 AIGC 技术在学术写作中的广泛应用，研究者在利用此类工具辅助写作的过程中，面临多重合规性风险，既包括学术不端的隐蔽性增强，也可能涉及对他人知识产权的侵犯，甚至因模型训练或交互记录导致个人隐私泄露等侵权问题。

3.1. 学术剽窃的合法性风险

在现行学术规范体系中，学术剽窃通常被界定为采用不当手段，窃取他人的观点、数据、图像、研究方法、文字表述等并以自己名义发表的行为。传统学术写作强调原创性思维与逻辑建构，研究者虽然可以合理接受导师或同行的指导建议，但写作主体仍应依托个人的独立构思完成。而随着 AIGC 在学术

写作中的广泛运用,学术写作面临前所未有的挑战。以 ChatGPT 等为代表的 AIGC 是一类基于 Transformer 神经网络架构,通过大规模数据训练与深度学习算法,具备理解自然语言并生成高质量文本的能力[9],其文本生成效率之高、语言结构之完备,使其在写作辅助中易被误用甚至滥用。在实践中,AIGC 的滥用正在催生若干新型学术不端行为。其一,直接替代性创作。研究者通过输入简单提示,由 AI 生成结构完整、逻辑自洽的学术文本,若将此类文本据为己有并署名发表,明显缺乏独立思维与实质创作,构成典型的生成式剽窃。其二,技术性“降重”行为。AI 工具常被用于对重复率偏高的文本进行“语义变换”或“语序重组”,实质未改变其对原始文献的依赖,而是掩盖了抄袭行为的表层特征,构成一种更为隐蔽的学术剽窃。此外,在学位论文或科研项目申报中,若研究者大量引入 AI 生成内容以规避查重机制,虽形式上呈现“新文本”,但实质上破坏了研究的原创性与真实性要求,严重违背学术诚信原则。

上述行为的合法性风险在于,AIGC 生成内容虽然在语言形式上表现为“新创作”,但其实质仍是对已有知识资源的高度重组与模仿。若研究者未对该等内容进行有效识别、核查与归属说明,即可能侵犯原始信息的著作权,并构成形式多样、难以识别的“技术性学术剽窃”,严重冲击既有的科研伦理和评价体系。同时,AIGC 系统在交互过程中通常需访问用户输入的信息以优化其生成结果。缺乏透明度和隐私保护可能导致数据被过度使用或用于不当目的,违反《中华人民共和国个人信息保护法》的合法性、透明性原则[10]。这一点也引发了关于研究者数据控制权和 AI 模型内容来源合法性的伦理争议。

3.2. 内容伪造和篡改的合法性风险

在学术出版领域,根据《学术出版规范》相关定义,学术伪造指通过编造或虚构数据、事实等手段获取学术成果;学术篡改则是故意修改或扭曲数据与事实,使其失去原有真实性。在 AIGC 时代,这两类传统学术不端行为呈现出新的技术形态和隐蔽特征。尤其是在基于大模型算法自动生成文本的过程中,人工智能“幻觉”现象(AI hallucination)极易导致对事实性信息的虚构或篡改,具体表现包括:生成不存在的文献引用;错误归属文献作者或出版年份;对事实数据的逻辑性篡改或统计性扭曲。这类虚假内容表面逻辑自洽、语言流畅,极具“迷惑性”,但本质上违背了学术写作中真实性、可验证性与可追溯性的基本规范。AIGC 系统在生成过程中,为了满足用户输入的逻辑预期,常在无依据或缺乏语料支持的情况下“自圆其说”,通过数据拼贴与推理链构造生成看似合理但实则失真的内容。相关研究指出,这一问题的产生与训练数据来源偏差、语料库真实性不一致、模型语义结构约束不足等多重技术因素密切相关[11]。此外,模型对人机交互中输入语义的误读或过度泛化,也可能触发对事实的任意重构与“合理化解释”,导致生成内容脱离客观数据基础。这些问题在学术论文撰写中尤其危险,可能误导研究方向、影响实证分析,最终损害研究的科学性与公信力。

更值得警惕的是,AIGC 在参与学位论文写作过程中,还可能侵犯学位申请人的基本权利,主要体现在对知情权与选择权的压制。一方面,由于当前大模型训练机制普遍依赖“算法黑箱”模式,训练语料集的组成、来源、筛选过程高度不透明,使用者难以了解 AI 生成内容所依赖的数据依据,进而丧失对所用信息来源的基本知情权。另一方面,“算法推荐”机制自动主导训练与生成过程,剥夺了使用者对参考文本、写作路径及推理逻辑的自主选择与控制能力,使其在不自觉中接受了模型内嵌偏见与预设立场,从而形成对选择权的结构性限制。从长远来看,人工智能参与学位论文创作可能导致学位论文的立场偏差或内容失真,进而在学位论文的实践转化及应用时产生更多的权利侵害隐忧。算法与数据是人工智能的技术基础[12],然而,由于人工智能并不具备自我意识[13],故其既无法对人为设置的算法偏见予以抗拒,也无法对其语料库中真假并存的数据从源头开展真伪审查。算法的偏见和数据的失真进一步导致人工智能生成的学位论文在立场与内容方面可能出现歧视或错误判断。这是导致权利侵害隐忧产生的技术根源。

3.3. 不当署名的合法性风险

不当署名是指学术成果的署名和顺序违背了实质学术成果的贡献规则，错配署名和资格。AIGC作为一种基于大规模预训练模型的技术，其核心逻辑在于通过对海量语料的学习，从中提取语言结构和知识表达的规律，再依据用户提示生成新的语言内容。尽管表面上产出的是“新”文本，实则是对既有知识与语言数据的重组与再现。因此，其在学术写作中的使用引发了署名合理性与知识产权合规性的双重争议。首先，AIGC在学位论文写作中的参与，可能导致对他人知识产权的侵犯。AIGC系统依赖训练数据集实现语言模型构建，而训练语料中往往包含大量受《著作权法》保护的文字作品、研究成果甚至学位论文文本。根据我国《著作权法》规定，凡具备“独创性”并由权利人独立完成、体现最低限度创造性的表达，均属于版权保护的客体范畴。因此，即便训练过程中未直接引用原始作品，一旦其生成内容本质上源自受保护文本，且未取得权利人授权，便可能构成侵权。考虑到现有著作权法对“独创性”的保护门槛较低，AIGC训练机制的合法性及其输出内容的权属问题应予以严肃审视。其次，AIGC的应用亦可能对学位申请人的隐私权构成威胁。AIGC模型构建所依赖的大规模语料库主要通过网络爬虫自动抓取信息或与用户互动获得语料。在用户与人工智能系统交互过程中，若不慎泄露个人敏感信息，如姓名、研究计划、论文草稿等，该信息可能被模型记录并在未来其他交互场景中被“复现”，从而引发隐私泄露风险。当前缺乏明确法律规范对AIGC系统中的信息采集与使用行为进行限制，使得隐私权保护面临现实挑战。

有学者指出，AIGC的文本生成行为本质上是一种“机械性创作”，其“创作能力”依赖于算法规则与大数据分析，缺乏真正的主观能动性与创造思维。正因为AIGC不具备人类作者应有的主体性与创新性，其生成内容难以满足署名所要求的“实质性贡献”标准。因此，将AIGC产出直接列入作者署名，不仅构成“学术不当署名”行为，也掩盖了可能存在的知识产权侵权和伦理问题。

4. 学术写作中生成性人工智能合法边界的探寻

4.1. 学术诚信与责任主体的界定

首先，使用者对AI辅助生成内容负有连带责任。《高等学校预防与处理学术不端行为办法》明确规定，未声明使用AI工具或剽窃AI生成文本的行为均构成学术不端。因此，研究者在利用AIGC提升写作效率的同时，必须对最终学术成果承担完全的学术与法律责任。换言之，任何AI生成的内容在提交前都应由作者本人仔细核查，并且在论文的致谢、方法说明或脚注中清晰标注其来源与用途；否则，一旦发生抄袭或数据造假等违规情形，研究者将无法以“工具使用”之名推脱责任。

与此同时，高校与学术期刊作为学术共同体的监管主体，也应履行严格的审查义务。一方面高校应建立分级问责机制，确保责任落实。在AIGC生成内容的过程中，涉及多方主体，如学生、导师、技术平台等。若发生学术不端或技术问题，学校应能追溯责任链。例如，若研究生生成不当内容，学校应明确责任来源，包括平台技术、学生行为和导师的指导责任。通过细化每个环节的责任，高校能迅速查明责任，确保合规性。另一方面，机构和期刊需将AI使用规范制度化，纳入投稿、答辩和评审流程。例如，可通过设置“AI使用声明表”要求作者详列所用工具名称、使用环节(如写作初稿、语言润色、数据可视化等)及人工审校方式。与此同时，应设立“AI审查小组”，对声明内容进行分级评估与抽查审核。该审查小组应由伦理审查委员会成员、领域专家与信息技术人员共同组成，确保兼顾学术专业性与技术判断力。其主要职责包括：(1) 审阅AI使用声明，判断是否存在对研究过程关键环节的过度依赖；(2) 开展样本论文溯源分析，通过文本检测工具识别潜在AI生成内容与原始语料重合度；(3) 就AI介入程度是否合理出具审核意见，并提出修改建议或驳回意见；(4) 定期评估现有AI使用规范执行情况，向期刊或

学校提出制度完善建议。

4.2. 著作权保护的法定边界

根据《著作权法》规定,凡具有“独创性”且已固定于特定载体的作品,均享有复制权、改编权、信息网络传播权等法定权能。因此,若将受著作权法保护的学术论文、图表或研究报告等文本在未经权利人授权的情况下直接纳入 AIGC 平台用于模型预训练或微调,就可能构成对上述专有权利的侵害。基于此,AIGC 服务提供者必须严格遵守以下合规要求:一是取得训练语料中所有受保护内容的合法授权;二是仅使用公共领域作品或符合开源许可的文本资源;三是于平台声明中明确列示语料来源及授权情况,以便使用者在调用前予以核验。与此同时,研究者在选择或调用 AIGC 平台时,也应主动审查其语料合法性承诺,确保所使用的训练数据不含未经许可的受保护作品,以免自身因“间接侵权”而承担连带法律责任。

AIGC 生成的文本本质上是一种基于既有语料的“重组性表达”,即模型通过算法将原有片段、句式及概念进行再组合,而非通过人类独立创意产生全新内容。根据我国《著作权法》及相关司法实践,著作权主体仅限自然人、法人或其他组织,且作品须体现“作者的智力创造成果”才能享有著作权。由此可见,纯粹由 AI 自动生成的文本不具备法定意义上的“独创性”,也不享有作者身份或相应著作权。然而,若研究者在 AI 输出之上进行实质性地改写、补充和再创作,并以自己的智力投入赋予作品新的个性化特征,则该再创作成果可由研究者本人以自然人身份申请著作权。反之,未满足“人类独创性”要件的 AI 输出内容,应视为公共领域素材,不得被研究者作为可申请著作权的原创成果。

4.3. 个人信息与隐私权的法律防线

首先,根据《中华人民共和国个人信息保护法》,在处理“敏感个人信息”及一般“个人信息”时,处理者必须遵循合法性、正当性、必要性原则,并在采集、存储或跨境传输前,明确告知信息主体处理目的、方式与范围,且获得其明示同意。因此,任何 AI 平台在未征得用户许可的情况下收集学术素材、研究数据或个人身份信息,均属违法。为此,平台运营者应当:一是提供“隐私模式”或“对话不保留”选项,允许研究者自主决定是否保留交互记录;二是在用户协议与隐私政策中,以简明易懂的语言说明对话数据的保存期限、使用目的及访问权限;三是针对涉密研究成果或高度敏感信息,实施差异化存储,并通过加密传输与访问控制机制,防止未经授权访问或数据泄露。通过上述措施,平台既能尊重用户的知情权与选择权,又能在技术层面构筑严密的隐私保护屏障。

其次,《AIGC 服务管理暂行办法》规定,AI 服务提供者须对用户交互内容进行分类分级管理,且仅能在用户明确同意的原业务场景内使用交互数据,严禁将其用于其他商业或非商业用途。为确保合规,平台应当:一方面,对不同类型的对话数据(如普通学术讨论与敏感科研数据)实施分级存储与权限管理,仅在特定用途下调用;另一方面,在对话界面或产品文档中清晰标示“本次对话数据是否用于模型优化”“数据保存时限”及“数据访问主体”,并提供随时查询与删除历史记录的功能。同时,研究者在涉及核心研究数据时,应主动选择或请求关闭训练回馈功能,以防止交互内容被二次利用;此外,还需认真阅读并遵循平台的合规说明,合理设置交互参数,从而在满足科研需求的前提下,有效规避隐私泄露风险,构建安全可信的学术写作环境。

4.4. 平台提供者的合规责任

信息安全与技术可追溯性是平台合规的基础。首先,依据《网络安全法》与《数据安全法》相关规定,平台应建立完善的日志留存与审计机制,对用户每次生成请求的提示词、模型版本、生成结果及后续编辑操作等关键信息进行完整记录与加密存储;日志数据应以加密形式存储,并严格限制访问权限,

仅授权经安全审查通过的运维与审计人员在特定情境下(如发生侵权投诉、数据泄露事件)进行调阅。此举不仅有助于平台在争议发生时及时锁定技术源头,还能向司法与监管机构提供合规责任的客观证据,增强系统的审计可用性与责任可归性。

当出现用户或第三针对数据泄露、非法调用等安全事件的举报时,平台不仅要立即启动应急预案,封堵漏洞并修复风险,同时还应积极配合司法机关与监管部门调取相应日志、提供技术支撑,以履行法律义务并维护各方合法权益。与此同时,合规提示与风险告知也是保障使用者合法合规的重要环节。为此,平台在用户界面或 API 文档中应显著设置免责声明,明确指出“本系统生成内容仅供参考,最终学术与法律责任由用户自行承担”;此外,还应提供专门的合规指南或 FAQ,详细阐述《著作权法》《个人信息保护法》等法律法规对 AI 使用的基本要求,并对潜在的侵权风险、隐私泄露风险及数据安全风险做出分类说明,以使用户在不同场景下采取相应的合规操作。这一机制不仅有助于提升用户自我合规意识,也能够平台与用户之间形成责任边界的“事前告知”和“事后免责”的证据链,防范误用带来的法律责任转移风险。

5. 结语

随着 AIGC 技术的不断演进,其在学术写作中的介入愈发深入,从语言润色到内容构建,AIGC 逐渐改变了学术写作的传统生态格局。本文在系统梳理国内外政策规制与学术规范演进的基础上,围绕学术剽窃、内容伪造、不当署名、知识产权侵权与隐私泄露等合规风险,界定了 AIGC 在学术写作中的合理使用边界与合法责任体系。研究指出,AIGC 在提升写作效率与表达质量的同时,其“拟人化输出”机制掩盖了源数据依赖与主观创造力缺失的本质,易催生隐性学术不端行为,挑战现有评价与伦理制度。为此,应从三个层面加以规范:一是研究者需强化责任意识,对 AI 辅助内容承担审校与标注义务;二是高校与期刊应完善 AI 使用披露与审核机制,将其纳入科研伦理治理框架;三是平台方须建立技术可追溯机制、数据合规声明与用户提示体系,确保系统运行的透明性与合法性。总而言之,AIGC 并非学术写作的威胁,而是驱动教育范式变革的技术引擎。唯有在法律规制、伦理引导与技术治理协同发力下,方能推动其在学术场域中“合规、安全、有序”地发展,助力知识创造与学术生态的可持续繁荣。

参考文献

- [1] 生成式人工智能服务管理暂行办法[EB/OL]. 2023-07-10.
https://www.gov.cn/zhengce/zhengceku/202307/content_6891752.htm, 2024-07-15.
- [2] Chan, C.K.Y. (2023) A Comprehensive AI Policy Education Framework for University Teaching and Learning. *International Journal of Educational Technology in Higher Education*, 20, Article No. 38.
<https://doi.org/10.1186/s41239-023-00408-3>
- [3] Thorp, H.H. (2023) ChatGPT Is Fun, But Not an Author. *Science*, 379, 313-313. <https://doi.org/10.1126/science.adg7879>
- [4] 负责任研究行为规范指引(2023) [EB/OL]. 2024-06-19.
https://www.most.gov.cn/kjbgz/202312/t20231221_189240.html, 2025-04-17.
- [5] Eacersall, D., Pretorius, L., Smirnov, I., et al. (2024) Navigating Ethical Challenges in Generative AI-Enhanced Research: The Ethical Framework for Responsible Generative AI Use. <https://arxiv.org/abs/2501.09021>
- [6] Bjelobaba, S., et al. (2024) Research Integrity and GenAI: A Systematic Analysis of Ethical Challenges Across Research Phases. <https://arxiv.org/abs/2412.10134>
- [7] Lo, C.K. (2023) What Is the Impact of ChatGPT on Education? A Rapid Review of the Literature. *Education Sciences*, 13, Article 410. <https://doi.org/10.3390/educsci13040410>
- [8] Rahman, M., Terano, H.J.R., Rahman, N., Salamzadeh, A. and Rahaman, S. (2023) ChatGPT and Academic Research: A Review and Recommendations Based on Practical Examples. *Journal of Education, Management and Development Studies*, 3, 1-12. <https://doi.org/10.52631/jemds.v3i1.175>
- [9] Brown, T., Mann, B., Ryder, N., et al. (2020) Language Models Are Few-Shot Learners. *Advances in Neural Information*

Processing Systems, **33**, 1877-1901.

- [10] 黄锬. 生成式 AI 对个人信息保护的挑战与风险规制[J]. 现代法学, 2024, 46(4): 101-115.
- [11] 韩博. ChatGPT 引发的人工智能内容生产传播风险[J]. 中国社会科学报, 2023(6): 34-36.
- [12] 贾开, 蒋余浩. 人工智能治理的三个基本问题: 技术逻辑、风险挑战与公共政策选择[J]. 中国行政管理, 2017(10): 40-45.
- [13] 赵磊磊, 吴小凡, 赵可云. 责任伦理: 教育人工智能风险治理的时代诉求[J]. 电化教育研究, 2022, 43(6): 32-38.