

人工智能工具医学使用的科学伦理与规范重构



白春学¹, 贾泽军¹, 杨达伟¹, 高承实^{2*}

1. 复旦大学附属中山医院, 上海 200032

2. 安徽栈谷科技有限公司, 池州 247100

[摘要] 随着人工智能技术的快速发展, AI 工具在医学研究与学术写作中展现出显著作用, 从数据处理、文献检索到论文写作、图表可视化和跨学科合作均实现效率提升。AI 不仅是辅助工具, 更逐渐成为科研伙伴, 在假说生成、实验设计和多模态数据分析中发挥创造性, 推动科研范式向“人机共研”转型。当前已出现多种应用场景, 如医学论文初稿生成、科研诚信检测、医学影像报告和科研图表的自动撰写。与此同时, AI 的广泛应用也对署名权、数据可追溯性、内容可信度及隐私保护提出挑战。国际指南如 ICMJE《医学期刊学术出版推荐规范》以及 Science、Nature 的公开声明均强调, AI 工具不得列为作者, 其使用需透明披露且责任归属人类研究者。为此, 亟需建立以透明披露、责任归属、质量可验证和合规开放为核心的规范体系, 以确保 AI 在医学科研中实现增效与创新的双重价值, 并为数字时代可持续科研生态奠定基础。

[关键词] 人工智能; 医学科研范式; 数据驱动科学; 生成式模型; 科研伦理

[中图分类号] R-3 **[文献标志码]** A

Reconstruction of scientific ethics and norms for the medical use of artificial intelligence tools

BAI Chunxue¹, JIA Zejun¹, YANG Dawei¹, GAO Chengshi^{2*}

1. Zhongshan Hospital, Fudan University, Shanghai 200032, China

2. Anhui Stack Alley Technology Co., Ltd, Chizhou 247100, Anhui, China

[Abstract] With the rapid development of artificial intelligence, AI tools have demonstrated significant value in medical research and academic writing, enhancing efficiency in data processing, literature retrieval, manuscript drafting, visualization, and interdisciplinary collaboration. Beyond serving as auxiliary tools, AI is increasingly becoming a research partner, contributing to hypothesis generation, experimental design, and multimodal data analysis, thereby fostering a paradigm shift toward “human - AI co-research”. Typical applications include rapid drafting of medical manuscripts, research integrity checks, and automated generation of imaging reports and scientific figures. However, the widespread adoption of AI also raises challenges concerning authorship, data traceability, content reliability, and privacy protection. International guidelines such as the ICMJE Recommendations and public statements from Science and Nature explicitly emphasize that AI tools cannot be listed as authors, that their use must be transparently disclosed, and that ultimate responsibility lies with human researchers. Therefore, it is urgent to establish a normative framework centered on transparency, accountability, verifiability, and compliant openness, ensuring that AI delivers both efficiency and innovation while laying the foundation for a sustainable research ecosystem in the digital era.

[Key Words] artificial intelligence; medical research paradigm; data-driven science; generative models; research ethics

近年来, 人工智能(artificial intelligence, AI)技术的迅速发展, 尤其是大语言模型、深度学习与多模态算法的广泛应用, 正在深刻改变医学研究与学术写作的方式。从医学影像智能诊断到临床数据挖掘, 从自动化文献综述生成到学术论文初稿撰写, AI 作为工具的能力不断突破传统科研的效率与边界。尤其是在 ChatGPT 等生成式语言模型的推动下, 科研人员在知识发现、论文写作与学术交流等环节中, 获得了前所未有的支持与便利, 甚至被

成为医生第二大脑。面对这一新兴力量, 医学研究和写作正站在一个新的历史关口, 即如何在充分发挥 AI 工具优势的同时, 确保学术的严谨性与科研的规范性?

传统意义上的学术规范与科研伦理, 是在以人工为中心的研究范式中逐步形成的。这一体系强调研究者对数据收集、分析与论文写作负有直接、不可替代的责任。然而, 在 AI 工具已经能够承担部分甚至大部分科研劳动的今天, 原有规范的适用

[收稿日期] 2025-06-23

[接受日期] 2025-09-15

[作者简介] 白春学, 博士, 教授, 主任医师. E-mail: bai.chunxue@zs-hospital.sh.cn

* 通信作者 (Corresponding author). Tel: 13838001036. E-mail: 13838001036@163.com

性正受到挑战。例如,研究人员是否可以在未披露的情况下使用 AI 生成文献综述? 如果论文的部分文本由 AI 起草,是否仍然符合署名与原创性的学术伦理? 面对这些问题,学术共同体在短期内往往采取“禁止”或“限制”策略,但这种基于旧有研究方式所延续的规范,很可能会阻碍新工具的合理使用,从而错失推动科学进步的重要契机。

医学研究与写作的核心目标是促进科学发展、推动其临床应用、改善人类健康。在这一根本目标的指引下,学术规范也应当与时俱进,主动吸纳能够提升科研效率与质量的新工具,而不是成为阻碍创新的桎梏。AI 工具不仅可以减轻研究者的重复性劳动,更能够在跨学科数据整合、临床证据快速生成、科研假说的形成与验证等方面发挥独特作用。如果能够在制度层面合理引导和规范 AI 的使用,它们将成为科研共同体的“增能工具”(augmentation tools),而不是学术伦理的威胁。

与此同时,我们必须正视 AI 工具可能带来的风险与挑战。一方面,生成模型仍然存在“幻觉”问题,其输出内容可能缺乏可验证性;另一方面,过度依赖 AI 可能引发学术诚信危机,模糊研究者与工具的责任边界。此外,医学研究具有高度的伦理敏感性,涉及患者隐私、临床数据安全与社会责任,AI 工具的应用必须严格嵌入医学伦理框架,避免技术滥用。因此,制定新的学术规范,不仅是技术问题,更是价值与责任的再平衡。

本文旨在全面审视 AI 工具在医学研究与论文写作中的应用现状,探讨其对科研方法论、成果呈现与学术规范的冲击,并提出新的伦理与规范框架,以实现效率、质量与责任的平衡。我们认为,学术共同体不应当回避 AI 工具的应用,而应以积极开放的姿态,推动规范重构,让人工智能成为医学科学的新型基础设施。这不仅是医学研究在数字时代的必然选择,也是实现科学可持续发展的关键路径。

1 AI 工具在医学研究与写作中的应用现状

AI 在医学中的应用已从最初的辅助性工具逐步发展为知识生产与科研范式转型的重要工具。医学是一门高度依赖数据与证据的学科,从影像学、临床记录到基因组学与科研文献,每一个环节都蕴含着潜在的规律。然而,长期以来,受制于计算能力与分析工具的不足,这些数据未能被充分挖掘。随着深度学习、自然语言处理(natural language processing, NLP)以及大语言模型(large language

model, LLMs)的突破,AI 已经能够对海量、异质化的医学数据进行系统化处理与知识抽取,由此推动科研流程在数据处理、成果写作以及协作教育等方面的深刻变革。

1.1 数据处理与知识发现 医学科研的核心在于通过数据揭示规律、发现新知。AI 技术的兴起,使数据处理与知识发现成为医学 AI 应用的前沿议题。

(1)医学影像诊断辅助。医学影像是 AI 技术率先落地的重要领域。早期研究集中在计算机辅助检测(computer aided detection, CAD),如乳腺 X 线片中标注可疑区域^[1]。随着卷积神经网络(convolutional neural networks, CNN)在 2012 年 ImageNet 挑战中的突破^[2],深度学习成为主流方法,并在肺结节检测、乳腺癌筛查、糖尿病视网膜病变识别等任务中展现出接近甚至超越专科医生的水平^[3]。

近年来,多模态影像(CT、MRI、超声)的结合联袂纵向数据建模使 AI 能够实现疾病动态演化的追踪。同时,研究重点从“替代医生诊断”转向“增强医生决策”。例如,Grad-CAM、注意力机制等可解释性方法揭示了模型的决策依据^[4],为医生提供额外参考。更重要的是,大规模影像数据的自动标注与分割逐渐成为临床科研的重要工具,推动疾病亚型探索与表型分层研究^[5]。

(2)大规模文献综述与知识图谱。医学研究高度依赖文献积累。传统的系统综述与 Meta 分析往往耗时数月至数年,难以应对文献爆炸式增长。早期的信息检索工具(如 PubMed 的关键词匹配)只能有限缓解^[6]。随着 NLP 与信息抽取技术的发展,大规模文献的自动筛选、分类与主题聚合成为可能。

近年来,LLMs 显著提升了文献处理能力。一方面,它们能基于数万篇文章生成研究趋势报告,捕捉跨领域联系^[7];另一方面,知识图谱逐渐成熟,能够将疾病、基因、药物关系结构化呈现,并支持药物再定位与新发疾病机制研究^[8]。由此,AI 已不再只是“加速综述工具”,而是可能成为科研发现的催化剂。

(3)临床数据挖掘与模式识别。电子病历(electronic health records, EHR)、可穿戴设备与多组学数据的积累,使临床研究进入大数据时代。早期数据挖掘依赖统计模型与浅层机器学习,用于风险预测或简单分类,但难以处理高维、非结构化与多模态数据。深度学习突破了这一限制,能够发现潜在模式与风险因子^[9]。

已有研究显示,基于 EHR 的预测模型在心衰、

败血症等风险预警中具有显著优势^[10]。同时,多组学数据的整合推动了精准医学的进展。例如,AI可通过模式识别优化临床试验入组,提升统计效力与临床相关性。更重要的是,AI模型还能通过聚类或潜变量建模发现新疾病亚群或机制性联系,从而推动疾病分类体系更新。

总体来看,AI在医学影像、文献综述与临床数据挖掘中,不仅提升了效率,更在生成新知识方面展现出潜力。它正逐步从“辅助工具”演化为“知识引擎”,为医学研究范式的转型奠定基础。

1.2 科研写作与成果呈现 科研写作是知识转化为学术成果的重要环节。随着AI技术的发展,学术写作的多个步骤开始被重塑。

(1)自动化摘要与文献综述生成。摘要与综述是科研写作中最耗时的部分。AI工具能够根据研究主题或初步数据自动生成研究背景与初步文献回顾,为研究者提供写作框架^[11]。这类工具并不替代学术判断,而是通过生成草稿减轻重复劳动,使研究者集中于学术价值与逻辑论证。已有研究表明,AI生成的综述在结构完整性与覆盖率上可与人工撰写相近,但仍需专家审校^[12]。

(2)学术语言优化与跨语种写作。英语是国际医学期刊的主要发表语言,长期以来对非母语研究者形成障碍。基于LLM的学术工具能够进行语法校正、风格优化甚至跨语种翻译。研究表明,这类工具能显著提高稿件被国际期刊接受的概率^[13]。在知识全球化的背景下,AI在语言优化中的作用不仅是“润色”,更是推动学术成果跨语种传播的重要途径。

(3)图表生成与可视化。数据可视化是科研表达的重要组成部分。传统图表工具在效率与复杂度方面存在局限,而AI能够根据数据自动生成统计图表,甚至实现复杂医学过程的三维可视化^[14]。近年来,AI驱动交互式可视化平台逐渐兴起,在临床研究报告和跨学科合作中显著提升了科研成果的可读性与影响力^[15]。

AI在科研写作环节的应用,核心价值并不在于替代研究者,而在于重塑学术表达的过程。它通过生成草稿、优化语言与可视化表达,降低了写作壁垒,推动了学术成果的全球传播与跨学科共享。

1.3 科研协作与教育 除数据与写作外,AI还在科研协作与教育领域展现潜力。医学科研往往依赖跨学科、多地域合作,而教育与培训则决定了科研共同体的可持续发展。

(1)智能助理与科研团队协作。科研团队日益

成为知识生产的主要单元^[16]。然而,跨地域、跨学科合作往往面临沟通与协调成本。AI工具作为“智能助理”,能够辅助会议安排、任务跟踪与实时文献检索。例如,自动生成会议纪要与行动清单的NLP系统已被用于科研团队管理^[17]。此外,AI驱动的知识图谱与推荐系统能揭示潜在跨学科联系,推动创新议题的产生^[18]。在新冠疫情等全球公共卫生危机中,这类工具加速了国际科研协作与知识共享^[19]。

(2)教育、培训与知识普及。医学教育正在经历由AI驱动的变革。自适应学习与智能辅导系统(intelligent tutoring system, ITS)能够根据学生表现动态调整教学进度与难度^[20]。在医学教育中,这种方法已用于病例推理、科研写作训练等场景^[21]。

在科研培训中,AI工具能模拟科研写作流程、自动批改研究设计,帮助青年学者快速掌握学术规范^[22]。在社会层面,AI驱动的开放科普平台能将复杂科研成果转化为通俗内容^[23],缩短科研与公众的知识鸿沟。

AI在科研协作与教育中的应用,推动了科研体系从“个体驱动”走向“智能协同”。它不仅提升了团队效率和教育质量,还扩展了知识传播的社会影响力。

综上所述,AI在医学科研的三个核心环节——数据处理与知识发现、科研写作与成果呈现、科研协作与教育——均展现出显著的学术价值和应用前景。在数据层面,AI实现了从影像识别到临床模式发现的知识提炼;在写作层面,它降低了科研表达与传播的门槛;在协作与教育层面,它推动了科研共同体的组织方式与人才培养模式的转型。简言之,AI正在成为医学知识生产的重要基础设施,为精准医学与转化医学的发展奠定坚实基础。

2 人工智能对科研规范与学术伦理的冲击

人工智能工具在医学研究与写作中展现出强大潜力,但与此同时,现有的科研规范和学术伦理体系正面临前所未有的挑战。传统的学术规则主要基于“人”在知识生产中的核心地位,而AI工具的广泛应用正在模糊研究者与工具的边界,从而引发了制度与实践的紧张关系。

2.1 现有学术规范的局限性 在深入分析AI带来的具体伦理挑战之前,有必要首先反思现有学术规范的生成逻辑及其内在局限。学术伦理体系并非一蹴而就,而是伴随着科研实践逐渐积累、形成的社会契约。其核心假设是,知识生产依赖于具有人

格与责任意识的研究者。然而,当 AI 工具在研究全过程中被广泛应用时,这一假设的适用性正在逐渐减弱。

(1)基于人工写作与传统研究方式形成的伦理体系。现有的学术规范与伦理框架,主要是在人工主导的研究与写作环境下逐渐确立的。署名权的界定、数据的管理规范,以及论文的原创性要求,均以“研究者是唯一知识创造主体”为前提。例如,《温哥华作者署名标准》明确要求作者必须对研究设计、数据分析和论文撰写承担实质性责任^[24]。然而,AI 工具的出现突破了这一假设,它们不仅能在写作中生成学术性文本,还能在研究过程中进行模式识别与数据推理,事实上参与了知识生产。由此引发的疑问包括,AI 生成的内容能否被认定为“原创”?若 AI 在研究设计中提出了重要假设,是否应被视为研究贡献?这些问题均超出了传统伦理框架的边界。

(2)对 AI 工具的使用缺乏明确规定。近年来,部分顶级期刊(如 *Science*、*Nature*)已发布声明,禁止将 AI 工具列为论文作者^[25]。然而,对其在研究与写作过程中的合理使用范围,学界尚未达成统一共识。一些学术机构采取保守排斥的态度,强调 AI 输出不可直接作为研究成果引用;另一些机构则主张有限容纳,允许在语言润色、翻译等环节使用 AI,但要求披露其作用。正如 Lund 等人^[26]所指出的,这种不一致性使得研究人员在使用 AI 工具时处于灰色地带,既担心违反规范,又难以放弃其高效性。换言之,制度供给滞后于实践需求,导致学术共同体在面对 AI 时缺乏明确的行为指南。

2.2 潜在风险 在理清学术规范的局限后,更为迫切的问题是识别 AI 在科研中的具体风险。这些风险不仅威胁到学术诚信,还可能动摇学术共同体赖以维系的信任机制。

(1)虚假数据和幻觉内容带来的学术诚信问题。生成式 AI 工具的“幻觉”现象,已成为学术界最为担忧的风险之一^[27]。所谓“幻觉”,是指 AI 在缺乏事实支撑的情况下生成虚假或错误内容,但表现形式却极为逼真。在医学研究中,这种现象可能导致引用不存在的文献、生成虚构的数据,甚至得出未经验证的结论。一旦缺乏严格的人工复核,研究成果的真实性与可重复性将受到严重威胁。与传统的学术不端(如抄袭、篡改数据)不同,AI 生成的虚假信息往往难以通过常规审查手段识别,从而对学术诚信构成系统性冲击。

(2)作者身份与署名伦理。AI 工具是否应被视

为论文的“作者”,是学术伦理面临的核心争议之一。根据现行的国际规范,作者必须具备承担法律与学术责任的能力,而 AI 显然不具备这一条件。因此,多数学术机构明确否认 AI 的作者资格。然而,完全忽视 AI 的作用也可能损害研究的透明度与可追溯性。学界正在探索折中方案,例如在论文的“方法”或“致谢”部分披露 AI 的具体作用^[28]。这表明,署名伦理的挑战并非在于是否承认 AI,而在于如何合理界定其“工具”属性与“贡献”性质。

(3)数据隐私与医学伦理冲突。医学研究高度依赖患者数据,而隐私保护始终是核心伦理议题。AI 工具在处理临床数据时,如果缺乏严格的数据脱敏与合规机制,可能导致敏感信息泄露。此外,AI 的“黑箱”特性意味着其决策过程难以解释,这与医学研究所强调的透明性和责任追溯相冲突^[29]。因此,如何在保障数据安全与患者隐私的同时,合理发挥 AI 的分析能力,成为医学伦理面临的又一重大难题。

2.3 制度滞后与实践冲突 在识别规范局限与潜在风险之后,必须正视一个现实问题,即 AI 在科研实践中的普及速度,已经远远超出了学术制度的响应速度。这种不协调不仅导致研究者的行为困惑,也削弱了学术共同体的制度权威。

(1)科研人员广泛使用与期刊、机构限制。大量调查^[30]显示,科研人员已在不同程度上依赖 AI 工具进行文献检索、语言润色和初步数据分析。然而,许多期刊和科研机构对此仍持保守态度,对 AI 使用缺乏系统性规范,甚至采取一刀切的禁止政策。这种制度与实践的错位,会导致研究者在实际操作中不得不“隐性使用”AI,从而形成事实上的规范失效与制度空转。

(2)制度更新滞后的风险。更为严重的是,学术规范的更新速度远远落后于 AI 技术的发展节奏。如果不能及时建立与之匹配的伦理与监管框架,既可能抑制 AI 的合理应用,也可能纵容无序使用带来的风险^[31]。更重要的是,制度滞后可能使学术共同体在国际规则制定中失去主动权,从而在全球学术竞争中处于不利地位。因此,制度与实践的错位不仅是技术应用的短期矛盾,更是学术治理的战略性挑战。

总体而言,人工智能工具对科研规范与学术伦理的冲击具有深远性和系统性。一方面,现有规范基于传统人工写作与研究方式建立,无法充分应对 AI 带来的原创性、署名与数据合规等问题;另一方面,生成式 AI 的幻觉风险、黑箱特性与隐私挑战,

进一步加剧了学术共同体的信任危机。在制度层面,科研人员的普遍使用与期刊、机构的保守限制之间存在显著张力,而规范更新的滞后则可能导致学术秩序的失衡。由此可见,医学研究这一伦理高度敏感的领域亟需建立一套既能保障科学严谨性,又能包容技术创新的全新学术规范与伦理框架。

3 新的科研工具使用伦理与规范框架

人工智能工具的广泛应用已成为医学研究与写作的重要趋势。面对新技术对学术规范与伦理体系的挑战,学术共同体不能仅停留在“禁止”与“限制”的被动应对层面,而应主动构建一套新的科研工具使用框架,使其既能保障科研质量与学术诚信,又能充分发挥 AI 的增效作用。本文提出四项核心原则,分别是工具透明化原则、责任归属原则、质量与可验证性原则、合规与开放原则。

3.1 工具透明化原则 在新型科研生态中,透明性被认为是重建信任的首要前提。缺乏披露不仅可能引发对学术成果真实性的质疑,还会加深学术共同体对 AI “隐性使用”的不安。因此,工具透明化原则应作为框架的基础性要求。

(1)明确披露使用范围与功能。科研人员在研究与写作过程中,应如实披露 AI 工具的使用情况,包括具体应用环节(如文献综述撰写、语言润色、数据模式识别)、所使用的工具类型(如 ChatGPT、Claude、Copilot 等),以及工具在成果形成中所扮演的角色。这一做法与学术研究对方法透明性的基本要求相一致,有助于提升成果的可追溯性^[32]。

(2)制度化披露机制。透明不仅依赖于个人自律,更需要制度保障。期刊与科研机构应建立统一的披露机制。例如,在论文投稿环节设立“AI 使用声明”模块,要求作者说明是否以及如何使用 AI 工具。通过制度化披露,可以有效避免因模糊使用引发的争议,同时为后续评估 AI 工具的实际学术影响提供系统性数据支撑^[33]。

3.2 责任归属原则 在强调透明的同时,更为关键的问题是如何界定责任。AI 工具虽能辅助研究,但其本质仍是工具,缺乏法律人格与责任意识。因此,责任归属原则应当成为学术伦理框架的核心支点。

论文作者要对最终内容负责任。无论 AI 工具在研究与写作中发挥何种作用,论文作者始终是科研成果的唯一责任主体。研究结论的科学性、数据的真实性、文字表述的准确性,都必须由作者亲自审查与确认。这一责任不可由 AI 工具分担或

替代。

AI 不具备署名资格。学界对 AI 是否应被列为作者的讨论已趋于共识,即 AI 工具不具备法律主体地位,无法承担学术不端或失实结论的责任,因此不应列入作者署名。然而,完全忽略 AI 的作用可能削弱成果透明度。因此,在方法学部分或致谢部分合理标注 AI 工具的辅助角色,既确保署名伦理的清晰性,又维护学术交流的开放性。

3.3 质量与可验证性原则 即便有了透明与责任机制,如果缺乏质量与可验证性的保障,科研成果仍将失去其学术价值。鉴于生成式 AI 存在“幻觉”和不稳定性,学术共同体必须构建可验证性机制。

(1)保证 AI 生成结果的可追溯性。医学研究对数据与结论的真实性有极高要求。科研人员在使用 AI 工具生成内容时,应确保结果均可追溯至可信来源。例如,AI 生成的文献综述必须逐一核实引用是否真实存在,并确认观点与原始研究相符。

(2)建立人工复核机制。生成式 AI 的输出不可直接等同于可靠证据。科研人员应将其作为“初稿”或“参考”,而非“最终版本”,并建立人工复核机制。在制度层面,期刊编辑与同行评审也应加强把关,确保论文不因未经验证的 AI 输出而损害学术严谨性。

(3)推动可验证的技术规范。从长远来看,开发者与学术界可合作推动 AI 工具引入可验证性设计。例如,自动生成参考来源、标记生成与编辑过程的“数字水印”。这些技术与规范结合的措施,有助于缓解“黑箱输出”带来的风险。

3.4 合规与开放原则 在强调透明、责任与质量之后,还必须正视更为宏观的制度问题,即如何在合规性与开放性之间寻求平衡。医学研究不仅涉及学术规范,还深受法律、隐私与伦理制度的制约。

(1)兼容医学伦理与数据隐私保护。医学研究高度依赖患者数据,因此 AI 工具的应用必须严格遵循隐私与伦理审查要求。临床数据在输入 AI 模型前应充分脱敏,且不得上传至未经审查的公共平台。同时,研究团队应优先选择本地化、安全可控的 AI 使用环境,以确保数据合规流转。

(2)鼓励国际化标准与共享规范。AI 工具的应用是全球性现象,规范的制定若仅局限于单一国家或地区,将导致制度割裂。国际学术机构如 ICMJE 与 WHO,可以在推动 AI 使用伦理与规范共识方面发挥关键作用。这种跨国框架不仅能增强学术共同体的信任,还能提升研究的可比性与协同力。

(3)平衡开放性与约束性

规范的制定既不能过于宽松,以致放任滥用;也不能过于严格,导致科研人员无法正当使用 AI。理想的做法是在开放与约束之间建立动态平衡,既承认 AI 工具的学术价值,鼓励其合理应用;又通过制度与伦理约束,防止其偏离科学研究的核心价值。

新的科研工具使用伦理与规范框架,不应仅是“补丁式”的修正,而应是一种系统性重构。通过确立工具透明化、责任归属、质量与可验证性、合规与开放四大原则,可以在保障学术诚信与科研质量的同时,最大限度发挥人工智能的增效作用。这一框架既回应了医学研究的现实需求,也为全球学术共同体在数字时代重塑学术规范与科学伦理提供了参考蓝图。

4 AI 赋能下的医学研究与写作新范式

人工智能工具的快速演化正在推动医学研究与写作进入全新的发展阶段。从最初的效率提升,到科研方法论的重构,再到学术共同体与出版体系的制度创新,AI 的介入不仅改变了研究的形式,更深刻影响了知识生产的逻辑。展望未来,AI 有望推动医学研究与写作迈向新的学术范式。

4.1 从效率工具到科研伙伴 当前,AI 工具多被视作“辅助性”工具,用于数据处理、语言润色或图表生成。但随着技术的进步,AI 将逐渐从效率工具演变为科研伙伴,能够在假说生成、研究设计与跨学科整合中发挥创造性作用。例如,AI 可以基于大规模医学数据库,快速提出潜在的因果关系假设,辅助研究者选择最优研究路径。

更进一步,AI 的推理能力将突破传统统计模型的局限,能够在多模态数据(影像、基因组、临床记录等)中发现复杂的非线性关系,为精准医学提供更为细致的研究视角。在这一过程中,研究者的角色将从单一的知识生产者,转向“问题提出者与判断者”,而 AI 则承担“探索与发现者”的功能,二者形成动态互动的人机共研模式。

4.2 可编程科学与智能实验设计的科研方法论再定义 AI 的嵌入将推动科研方法论发生根本性变革。传统科研强调人工制定实验设计与数据分析方案,而在 AI 赋能下,科研将走向“可编程科学”。研究人员可以通过自然语言或代码指令,调用 AI 模型自动化完成实验流程,从假设生成到数据收集、从建模分析到结果解释,形成高度智能化的“端到端科研链”。

这一变革带来的最大优势在于动态迭代与实

时优化。不同于传统实验设计的固定性,AI 驱动的实验可以在过程中不断修正假设和参数配置,实现“边研究边优化”。例如,在临床试验中,AI 能够根据实时数据反馈动态调整受试者分组策略,从而缩短实验周期,提高试验效率与结论可靠性。这种“智能实验”模式,不仅提升了科研效率,还可能重塑科学的基本范式,即从“人类驱动”走向“人机协同驱动”的知识生产逻辑。

4.3 期刊与学术共同体的规范创新 随着 AI 的深度应用,学术出版与共同体运作也将迎来结构性革新。

(1)出版流程智能化。未来,医学期刊的出版流程将因 AI 的介入而大幅革新。AI 可以在稿件初审中自动完成语言校对、格式核查、相似度检测,甚至进行初步的逻辑一致性评估。这不仅缩短了出版周期,也能将人力资源释放到更高价值的内容评估环节,从而提升整体出版效率。

(2)审稿机制革新。AI 也将改变同行评审机制。未来可能出现“人机联合审稿”模式。AI 工具先对论文的事实准确性、数据可验证性进行快速审查,随后由专家学者对研究价值与创新性进行深度评议。这种双重机制不仅能够减轻审稿压力,还能提升评审的客观性与一致性。更重要的是,通过大规模积累 AI 辅助审稿数据,还能推动形成客观化的“科研质量指标”,减少传统同行评审中的主观偏差。

4.4 元宇宙医学语境下的 AI 学术生态 在进一步展望中,AI 的学术功能不应局限于当下的研究与出版体系,更应该将其延伸至更具沉浸感和协作性的未来科研空间中。随着元宇宙概念与医学研究的交融,AI 将作为底层驱动力,支撑起跨空间、跨学科、跨机构的科研互动与知识共建。

(1)跨空间与沉浸式科研协作。在元宇宙医学的发展框架下,AI 工具的作用将进一步扩展至跨空间、跨学科的知识交互与协作中。借助沉浸式环境与虚拟科研空间,研究人员可以在元宇宙中共享实时实验数据,利用 AI 进行智能化模拟实验与可视化呈现,实现“虚拟临床试验”与“数字孪生患者”的学术探索。

(2)智能化虚拟科研基础设施。AI 将成为元宇宙科研生态的“智能中枢”。例如,在虚拟实验室中,AI 可以自动化协调资源调度、实验流程管理与数据分析;在全球科研网络中,AI 可作为翻译、对话与推理的中介,使跨语言、跨文化的科学交流无缝进行。这一趋势意味着,AI 不仅是研究过程中的

“助手”,更是未来学术生态的制度性与技术性基础设施。

总体来看,人工智能正在推动医学研究与写作实现三重转型,即从效率驱动到伙伴协作,从传统方法论到智能化科学,从静态规范到动态共建,其影响不仅限于提高科研生产力,更将深刻地体现在对学术规范、知识逻辑与学术生态的重构上。在元宇宙医学的语境下,AI 不仅是技术手段,更是学术共同体的新型基础设施。

未来,如何在促进科研创新的同时确保学术诚信、伦理合规与社会责任,将成为学术界与产业界共同面临的重大课题。只有通过科学合理的制度设计与跨学科的共同治理,人机共研的学术生态才能真正加速人类对生命健康的认知与实践。

5 结 语

人工智能工具的迅速发展普及,正以前所未有的深度与广度重塑医学研究与学术写作的生态。从数据处理、文献检索到知识发现,从科研写作、图表可视化到学术交流与合作,AI 的介入不仅显著提高了效率,更在研究设计、假说生成及跨学科整合中发挥了潜在的创造性作用。可以说,AI 不仅是研究者的“智能助手”,更正在成为医学知识生产的“伙伴”,甚至成为医生的第二大脑,推动科研从传统的人力密集模式向“人机共研”的协作化、智能化新范式转型。

然而,这一技术变革同时也带来了学术伦理、数据安全与制度规范的深层挑战。现有学术规范体系过度依赖人工研究与写作假设,未能充分覆盖 AI 的使用场景,导致署名权、数据可追溯性、内容可信度及患者隐私保护等方面出现潜在冲突。制度滞后与实践高速发展的不平衡,既可能引发学术诚信风险,也可能阻碍科研创新的合理开展。

因此,学术共同体必须主动建构一套新的科研伦理与规范框架。核心原则应涵盖透明披露 AI 使用、明确责任归属、保证结果质量与可验证性、遵循合规与开放标准。在此基础上,在确保科研严谨性与社会责任的前提下,AI 工具的应用才能够进一步实现科学发现的加速、跨学科合作的拓展和医学研究模式的智能化升级。长远来看,AI 将成为医学科研与知识生产的基础设施,与大数据、云计算和其他数字技术共同构建数字时代的可持续科研生态。

伦理声明 无。

利益冲突 所有作者声明不存在利益冲突。

作者贡献 白春学:选题、修改论文;贾泽军、杨达伟:参与讨论;高承实:选题、撰写、修改论文。

参考文献

- [1] DOI K. Computer-aided diagnosis in medical imaging: historical review, current status and future potential [J]. Comput Med Imaging Graph, 2007, 31(4/5): 198–211.
- [2] KRIZHEVSKY A, SUTSKEVER I, HINTON G E. ImageNet classification with deep convolutional neural networks [C]// Advances in Neural Information Processing Systems 25. 2012, 1097–1105.
- [3] ESTEVA A, KUPREL B, NOVOA R A, et al. Dermatologist-level classification of skin cancer with deep neural networks [J]. Nature, 2017, 542(7639): 115–118.
- [4] SELVARAJU R R, COGSWELL M, DAS A, et al. Grad-CAM: visual explanations from deep networks via gradient-based localization [C]//2017 IEEE International Conference on Computer Vision (ICCV). October 22–29, 2017, Venice, Italy. IEEE, 2017: 618–626.
- [5] LITJENS G, KOOI T, BEJNORDI B E, et al. A survey on deep learning in medical image analysis [J]. Med Image Anal, 2017, 42: 60–88.
- [6] MOHER D, LIBERATI A, TETZLAFF J, et al. 系统评价与荟萃分析优先报告条目 (PRISMA) 声明 [J]. 中国循证医学杂志, 2009, 9(9): 1085–1091.
- [7] LEE B, LEE S, KIM D, et al. Large language models for systematic review automation: evidence synthesis at scale [J]. Journal of Medical Internet Research, 2023, 25: e48625.
- [8] PERCHA B, ALTMAN R B. A global network of biomedical relationships derived from text [J]. Bioinformatics, 2018, 34(15): 2614–2624.
- [9] MIOTTO R, LI L, KIDD B A, et al. Deep patient: an unsupervised representation to predict the future of patients from the electronic health records [J]. Sci Rep, 2016, 6: 26094.
- [10] SHICKEL B, TIGHE P J, BIHORAC A, et al. Deep EHR: a survey of recent advances in deep learning techniques for electronic health record (EHR) analysis [J]. IEEE J Biomed Health Inform, 2018, 22(5): 1589–1604.
- [11] ELHADAD N. AI-assisted drafting of literature reviews: benefits and limitations [J]. Journal of the American Medical Informatics Association, 2019, 26(8–9): 785–787.
- [12] HOPE T, CLARK J, BENDER E M, et al. Can large language models accurately summarize systematic reviews? A human evaluation [J]. Journal of the American Medical Informatics Association, 2022, 29(11): 1923–1932.
- [13] KIM Y, LEE M, KIM D, et al. (2023). Towards Explainable AI Writing Assistants for Non-native English Speakers. arXiv.
- [14] RAUBER P E, FADEL S G, FALCÃO A X, et al. Visualizing the hidden activity of artificial neural networks [J]. IEEE Trans Vis Comput Graph, 2017, 23(1): 101–110.
- [15] ZHANG X, WANG Y, LIU C. AI-driven interactive visualization platforms in clinical research reporting: a

- systematic review[J]. *Journal of Biomedical Informatics*, 2023, 140: 104339.
- [16] WUCHTY S, JONES B F, UZZI B. The increasing dominance of teams in production of knowledge [J]. *Science*, 2007, 316 (5827): 1036–1039.
- [17] SEE A, LIU P J, MANNING C D. Get to the point: summarization with pointer-generator networks [C]//*Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*. Vancouver, Canada. Stroudsburg, PA, USA: ACL, 2017.
- [18] PERONI S, SHOTTON D. OpenCitations, an infrastructure organization for open scholarship [J]. *Quant Sci Stud*, 2020, 1 (1): 428–444.
- [19] NIAZI M K K, SIDDIQUI S A. COVID-19 research collaboration networks: a 50-country analysis [J]. *Scientometrics*, 2020, 125(3): 2557–2576.
- [20] KOEDINGER K R, CORBETT A T, PERFETTI C. The knowledge-learning-instruction framework: bridging the science-practice chasm to enhance robust student learning [J]. *Cogn Sci*, 2012, 36(5): 757–798.
- [21] CHAN L, GOH P S, WONG M L. Adaptive learning and intelligent tutoring systems in medical education: a systematic review [J]. *Medical Teacher*, 2022, 44(4): 392–401.
- [22] ZHANG Y, WANG X, LIU C. AI-driven scientific writing simulators for early-career researchers: development and evaluation [J]. *Journal of Medical Internet Research*, 2023, 25: e42851.
- [23] MINH T N, HOANG T Q. AI-powered open science communication platforms: bridging the gap between research and public understanding during COVID-19 [J]. *Health Communication*, 2021, 36(14): 2145–2154.
- [24] International Committee of Medical Journal Editors. Recommendations for the Conduct, Reporting, Editing, and Publication of Scholarly Work in Medical Journals (ICMJE Recommendations 2023) [S/OL]. (2023-12-18) [2025-09-02]. <https://www.icmje.org/recommendations/>.
- [25] Nature. AI tools such as ChatGPT must not be listed as authors on research papers [N]. *Nature News*, (2023-01-27) [2025-09-02]. <https://www.nature.com/articles/d41586-023-00107-4>.
- [26] LUND B D, WANG T, MANNURU N R, et al. ChatGPT and a new academic reality: AI-assisted research and the erosion of authorial integrity [J]. *Journal of Academic Librarianship*, 2023, 49(4): 102589.
- [27] JI Z W, LEE N, FRIESKE R, et al. Survey of hallucination in natural language generation [J]. *ACM Comput Surv*, 2023, 55 (12): 1–38.
- [28] HOLDEN T H. ChatGPT is fun, but not an author [J]. *Science*, 2023, 379(6630): 313.
- [29] PRICE W N, COHEN I G. Privacy in the age of medical big data [J]. *Nat Med*, 2019, 25: 37–43.
- [30] VAN DIS E A M, BOLLEN J, ZUIDEMA W, et al. ChatGPT: five priorities for research [J]. *Nature*, 2023, 614 (7947): 224–226.
- [31] FLORIDI L. The governance of AI: ethical and policy challenges [J]. *Nature Machine Intelligence*, 2023, 5 (8): 827–829.
- [32] ANDERSON K, CHEN L, ZHANG Y, et al. Transparent reporting of AI assistance in scientific writing: a framework for disclosure [J]. *Research Integrity and Peer Review*, 2023 (8): 12.
- [33] Nature Editorial. AI and the future of writing: transparency is key [J]. *Nature*, 2023, 618(7962): 214–215.

引用本文

白春学, 贾泽军, 杨达伟, 等. 人工智能工具医学使用的科学伦理与规范重构 [J]. *元宇宙医学*, 2025, 2(3): 48–55.

BAI C X, JIA Z J, YANG D W, et al. Reconstruction of scientific ethics and norms for the medical use of artificial intelligence tools [J]. *Metaverse Med*, 2025, 2(3): 48–55.