

CERI

长江教育研究院

Changjiang Education Research Institute

教育智库
Thinktanks

2026年05月刊

总第93期



P01 卷首语：

潘教峰：人工智能赋能的本质认识：数据、知识与系统的三重整合

目录 Contents

卷首语

- 01 人工智能赋能的本质认识：数据、知识与系统的三重整合 潘教峰

专题聚焦：教育强国建设核心议题探索

- 19 技术、制度与思想：生成式人工智能在教育领域中的应用的演进逻辑 周洪宇
- 31 推动人工智能融入教育全要素全过程全场景 杨宗凯
- 35 生成式人工智能时代高等教育变革：本体挑战、异化风险与体系重构 邬志辉

专题研究：教育强国建设实践探索

- 49 联合国教科文组织人工智能教育立场的阐释与实践启示 王小飞
- 61 欧美国家人工智能治理话语的计算分析：主题特征与演化趋势 李佐文
- 73 生成式人工智能时代教育创新范式的演进路向 龙宝新



欢迎与我们互动

目录 Contents

院内动态

- 84 AI 赋能教育创新 传承行知思想 —— 苏州市教育学会、陶行知研究会专项培训圆满落幕
- 91 周洪宇院长受邀出席《智慧教育》创刊暨资源中心期刊编委会成立座谈会
- 94 周洪宇院长受邀参加教育强国建设重大问题研究与中国教育学体系建构学术研讨会
- 95 周洪宇院长受邀出席第三届教育改革与发展论坛暨《河北师范大学学报（教育科学版）》新一届编委会会议
- 97 周洪宇院长受邀出席红色教育数智记忆实验室建设与发展研讨会
- 103 周洪宇院长受邀出席“求是教育论坛”

征稿通知

- 106 《教育治理研究》的征稿通知



欢迎与我们互动

人工智能赋能的本质认识：数据、 知识与系统的三重整合

来源 | 《科研管理》2026 年 01 期



潘教峰

第十四届全国人大代表
中国科学院大学公共政策与管理学院院长
中国科学院科技战略咨询研究院院长

人工智能作为引领新一轮科技革命和产业变革的战略性技术深刻改变人类生产生活方式。OpenAI 的首席执行官山姆·奥特曼曾指出人工智能革命是继农耕时代、工业时代、计算机时代之后的第四次革命有望推动社会实现共同富裕。回顾技术革命历史人类社会的

重大变革是通过机械力、智慧力、繁殖力的不断外化逐渐形成以人与自然、人与机器和谐相处为中心的社会新形态。从蒸汽机的出现到内燃机和电力的发明机器对人力的替代重构社会生产方式是“机械力”外化的象征而电子计算机的应用将人类认知功能外化为机器可执行的程序是关乎“智慧力”的革命。人工智能则通过对人类执行层和决策层的双重替代在延伸“机械力”与“智慧力”的基础上展现出颠覆式的“创造力”。

顺应世界科技发展大势 2024 年 3 月我国首次将“人工智能+”行动列入《政府工作报告》标志着我国人工智能进入全面赋能新阶段智能制造、智能驾驶、智慧城市、智慧医疗等新业态、新模式快速涌现。人工智能在赋能过程中表现出高度“整合”的特征。从字面意思看整合是指“将不同的部分或元素组合在一起形成一个统一的整体”在计算机领域整合是“通过互联网技术实现跨地域、跨组织的资源协调与重组”。在人工智能赋能的语境中整合不仅是技术的简单应用更是对组织数据、知识的融合创新是与场景需求的深度融合。邹起浩与任保平将新一代人工智能驱动的新型工业化视为一种“技术—组织—产业”的新范式指出实体经济和数字经济通过生产层面和生态层面的深度融合实现资源与供需的高效匹配。董雪艳等认为在人机融合的治理模式下 AI 发挥信息整合、联动协同、决策支持、指挥救援等方面的潜能结合人类经验价值判断将带来应急管理全过程的提质增效。文军与陈宇涵的研究指出人工智能自主“思考”与情感交互的增强正超越预设的“工具”界限通过驱动情感的识别、生成与传播促进社会工作中情感的精细化、专业化与多元化发展。人工智能通过虚拟协同系统构建与现实运行机制耦合正重塑物质流、信息流与价值流的组织逻辑推动形成一种新型的社会智能形态影响社会结构与发展范式的深刻变革。

在此背景下本文旨在从整合视角出发把握人工智能赋能的底层学科逻辑厘清人工智能赋能整合的内在机理与逻辑架构并从系统论

的角度明确当前人工智能赋能整合面临的风险挑战提出深化人工智能赋能整合的对策建议以更好地规制和使用人工智能。

1. 整合视角下的人工智能主要流派

人工智能作为一种具有高度整合能力的通用技术是典型的多学科、多领域融合的交叉学科其主要流派的发展无不蕴含对计算机科学、认知科学、神经科学、语言学、数学、控制论等众多底层基础学科的融合与重构。

1.1 符号主义下的“规则—逻辑”整合

符号主义作为人工智能初期探索和发展阶段的主流学派起源于数理推断和语言学主要依赖于符号处理和规则模拟系统把人类思维符号化表达并通过形式逻辑和规则组合来模拟人类的推理过程。

1956 年达特茅斯会议约翰·麦卡锡首次提出“人工智能”AI 概念拉开符号主义模拟人类智能探索的序幕。1956 年艾伦·纽厄尔和赫伯特·西蒙开发的第一个人工智能程序——“逻辑理论家”LT 证明了 38 条数学定理。1965 年第一个专家系统 DENDRAL 问世能够根据分子式和质谱数据推断化合物分子结构证实了专家系统在实际应用中的可行性。此后 DEC 计算机配置系统 XCON 的成功以及 IBM 超级电脑“深蓝”的胜利使得符号主义一时间声名大噪。

符号主义将“推理”转译为可编程符号系统建立了形式化语言与知识表达的转译和整合桥梁一定程度上赋予机器智能。但其静态知识库与单一推理结构在处理复杂、不确定性问题时缺乏灵活性和自适应能力限制了真正“智能”的实现。

1.2 行为主义下的“控制—环境”整合

行为主义最早起源于控制论强调将所有的行为都视为对外部刺激的反应在发展中逐步融入心理学、生物学与机器人学交叉探索智能体与环境的动态耦合。

20 世纪 40—50 年代控制论思潮的兴起倡导基于控制科学的理论、方法和技术模拟人类智能行为。1966 年美国斯坦福研究所 SRI 研发的第一台通用机器人 **Shakey** 能够通过自身行为“推理”解决简单的感知、运动规划和控制问题。20 世纪 80 年代“感知—行为”模式理论的形成标志着行为主义作为人工智能的一个重要学派开始确立。该理论强调获取环境信息以指导智能系统的行为决策。1989 年罗德尼·布鲁克斯制造的六足机器人 **Genghis** 采用基于“感知—动作”模式模拟昆虫行为的控制系统被视为新一代的“控制论动物”。

行为主义通过控制工程和认知科学耦合实现系统与环境的互动使 AI 真正“落地”到物理世界为自动驾驶、智能制造等场景奠定“环境交互”的系统模型。但高昂的试错成本、复杂场景适应性不足等缺陷导致行为主义发展滞后于符号主义和连接主义。

1.3 连接主义下的“神经—统计”整合

连接主义以脑科学为指导强调人脑工作原理的模拟。该学派试图通过融合神经生理学、物理模拟与数学模拟等多种学科方法构建类似人脑神经元节点网络的仿生结构以实现机器思维 and 智能学习。

1943 年首个人工神经元模型——MP 模型的提出开启了连接主义的篇章 1958 年弗兰克·罗森布拉特发明的“感知机” **Perceptron** 模型奠定神经网络的雏形。20 世纪 80 年代随着神经网络的复苏机器学习方法逐渐取代专家系统成为研究的核心方向。1986 年杰弗里·辛顿等人提出的 BP 算法以反向传播的方式实现机器的自我学习增强了系统的适应力。2006 年深度学习概念的提出通过设置更多的隐藏层可达数百个自动提取海量数据特征实现隐藏特征的深度挖掘展现出超越人类的智能整合能力。在此基础上生成对抗网络 **GANs** 与 **Transformer** 模型的提出引发生成式人工智能的爆发式增长。**BERT**、**GPT** 系列、**Gemini**、**Sora**、**DeepSeek** 等通过概率驱动的内容生成实现从“发现”到“创造”的跨越。

连接主义以脑科学为指导整合不同学科的知识、数据并通过深度

学习推动模型泛化能力的不断提升。但不容忽视的是当前人类对脑科学认知的不足也制约了连接主义的发展造成算法黑箱与不可解释性等局限这意味着未来人工智能的进一步发展仍离不开底层学科的突破。

1.4 三大流派的融合

作为自然科学、社会科学、技术科学交叉创新的成果虽然符号主义、行为主义、连接主义在机理和应用上各有侧重但人工智能的各流派也呈现出相互融合的趋势。例如行为主义与连接主义的融合可以使人工智能系统通过环境交互动态调整自身神经网络结构从而增强学习效率和适应能力。以自动驾驶场景为例将行为主义的强化学习算法与连接主义的深度学习模型结合可以使自动驾驶系统在复杂交通环境中更加高效地学习并进行行为优化提升安全性和可靠性。符号主义、行为主义与连接主义的全面融合将实现数据驱动的统计学习、知识主导的逻辑推理以及系统与环境交互的高度统一。这不仅将弥补单一流派存在的局限性如符号主义的静态性、行为主义的高试错成本、连接主义的算法黑箱等问题还将为实现人工智能在复杂、多变环境中的应用奠定坚实的科学基础从而推动更加全面的智能赋能与突破。

2 人工智能赋能的本质

2.1 人工智能赋能的整合内涵

人工智能以底层学科知识为支撑展现出强大的普适性、通用性广泛服务于经济社会各个领域的应用需求。其赋能本质在于依托“多元数据—五大能力—三类技术”的融合实现“数据—知识—系统”三重整合构建“感知—认知—执行”的赋能闭环。

数据整合是人工智能赋能的根基通过多模态、跨领域、跨时空数据采集与融合实现对物理世界状态的全面表征构建起对复杂环境的数字化映射。

知识整合是人工智能从数据中提炼规律与逻辑实现对问题抽象

的过程其背后是AI基于统计学习、逻辑推理等方法构建的“五大能力”关联识别能力挖掘变量之间的共现性因果推理能力揭示事物发展的内在机理矛盾发现能力解决传统理论和方法难以解决的问题收敛逼近能力寻求复杂系统的计算实现途径突变涌现能力突破常规式创新。这些能力推动了人工智能从感知走向认知从而使赋能经济社会各领域成为现实。

系统整合通过基础技术、功能技术、领域技术的综合集成推动从单点技术到系统工程的实现在虚拟世界中构建起对物理世界的映射与操控体系。通过虚拟世界中的设备、系统与流程实现对物理世界的控制推动“人—机—物”三元深度融合使人工智能优化现实世界运行效能成为可能。



图1 人工智能赋能整合的本质机理
Figure 1 Essential mechanism of AI-enabled integration

2.2 数据整合

数据是人工智能赋能实现系统性整合的关键基础。传统模式下模型处理的数据类型较为单一且数据多处于分散化、碎片化的状态难以提供决策所需的完整信息。人工智能通过整合多源异构数据实现了跨时空、跨领域、跨组织数据的相互支持与补充提供更加准确、全面、及时的信息支持。

人工智能实现多模态数据整合。多模态数据整合是指将来源不同、形式各异的数据类型进行融合包括图像、语音、文本、视频、

传感器信号等。人工智能通过构建统一的表征学习框架与深度神经网络结构使异构数据能够在同一语义空间中进行关联与解析。多模态数据整合能够利用不同模态数据间的互补性提高模型性能即通过捕捉多模态数据间的交互信息开展复杂系统建模的同时降低单一模态中低质量、错误数据与异常、缺失数据对模型构建的影响做出更具可靠性和鲁棒性的预测。多模态数据整合在自动驾驶、医疗等领域具有重大的应用价值特别在医疗领域传统医疗研究通常面临生理参数、生化指标等结构化数据难以与 X 光片、磁共振成像等非结构化数据整合的难题人工智能通过多模态融合技术有效打破结构化与非结构化数据之间的壁垒实现疾病状态的全方位刻画和精准识别。

人工智能实现跨时间数据整合。数据作为信息承载的重要媒介记录了事物的发展演进过程也蕴含着事物的内在机理。人工智能以数据为载体可实现历史域、现实域和未来域的融合贯通。历史域中历史隐性信息的显性转化为累积资源调用提供了可计算、可迭代的智能资产。现实域中全域数据抓取、实时融合解析与动态预警监测大幅提高数据覆盖广度与信息触达速度打破传统信息的滞后性实现关键指标的动态追踪。未来域上通过时序建模、时间序列分析、因果推理等方法挖掘历史与现实数据随时间推移所呈现出的规律性、周期性或突发性特征能够实现未来趋势预测。例如谷歌“深度思维”开发的名为 **GraphCast** 的人工智能天气预报模型采用近 40 年历史天气数据进行训练由此产生的算法能够在一台台式电脑上于一分钟内预测未来 10 天的全球天气且 90% 的指标明显优于当前的天气系统。

人工智能实现跨区域空间数据整合。空间数据是描述自然地理空间和人类活动空间中所包含的人、物体、事件的信息。通常来说空间数据具有空间位置信息、时间信息和属性信息。随着各类传感器和全球定位系统的广泛使用空天遥感、地图测绘、交通轨迹、手机信令、APP 签到等空间数据持续增长成为全域互联新的智能接口。与传统地理信息系统 GIS 相比 AI 赋能的空间数据整合不仅提高了对海量空间数据的自动感知与处理能力更通过深度学习、图神经网络、

空间—时间建模等先进方法实现了空间数据的结构化表达、语义理解与动态演化预测。高度整合的空间数据不仅能够为城市交通管理与能源规划提供基础支持更通过揭示人口关联的空间信息为全流程的防灾应急与传染病防疫提供决策依据。

人工智能技术实现跨领域、跨部门数据整合。开放共享的数据管理平台进一步打通全域数据池实现跨领域、跨部门数据的整合。这不仅能够促进不同领域、不同部门之间的信息流通实现跨界资源共享与统一调配经验知识的协同复用还能够降低重复劳动、激发群体智能从而极大地提升工作效率和决策的科学性。例如广东省应急管理厅成功构建“一网统管风险防控与应急指挥体系”该体系接入多达 27 个外部厅局以及 14 个应急厅内部机构涵盖 1171 类业务数据并提供 1372 类数据服务。这一体系的建立不仅实现对各类风险的全天候监控预警还能够在紧急情况下进行跨层级、跨区域的统一指挥调度极大地提升了应急管理的效率和水平。

2.3 知识整合

从人工智能发展历程及流派演化可见 AI 的跃迁源于对多学科知识的深度整合。而在数据的强力支撑下人工智能已广泛嵌入经济社会的各个领域从“优化工具”转变为“创新引擎”。它不仅提升了传统流程效率更通过主动参与设计生成、方案评估和科研运营等复杂过程系统性地加速了创新周期与知识演化。这背后的核心驱动机制体现在人工智能所具备的五大关键能力关联识别能力、因果推理能力、矛盾发现能力、收敛逼近能力、突变涌现能力。

关联识别能力揭示变量间的深层联系。基于统计学和概率论的基础假设人工智能能够通过分析海量数据利用机器学习算法识别变量间的相关性。这不仅是在原有数据上的简单聚类更通过非线性的特征识别挖掘数据内在的隐含关联展现出对表层关联和结构关联的深度分析。在表层关联中 Apriori 算法通过逐层搜索和剪枝策略揭示数据项间的频繁共现关系沃尔玛便是基于此从购物篮数据中挖掘出“啤

酒与尿布”的强关联规则。这些应用基于数据的表层关联或共现关系识别出潜在的相关性。而进一步的 DeepMind 发布的 AlphaFold 3 引入扩散模型进行多尺度特征学习并结合 Cryo—EM、NMR 等实验数据修正扩散过程偏差将蛋白质—配体相互作用的预测精确度提升至 76% 相较最优的传统方法高出约 50% 这实现了人工智能从“表层相关”走向“结构性关联”的飞跃。关联识别作为人工智能最基本的能力奠定了技术赋能的根基。

因果推理能力是超越相关性的机理推理。关联识别的揭示虽能捕捉事物潜在的相关性但也面临“虚假”相关的问题而因果推理则从深层次反映数据生成机制间的关系这种关系更具稳健性和泛化性。因而要把握事物发展的根本规律仍需从因果推理出发。图灵奖获得者朱迪亚·珀尔 Judea Pearl 认为人类通过三层认知阶梯来理解世界即观察关联干预行动以及反事实推演。干预行动和反事实推演都必须建立在因果推断之上因而因果推断的算法化是实现强人工智能的关键。对此珀尔提出的因果图清晰地表达了变量之间的因果结构并能够利用图的性质来推导因果效应奠定了构建因果人工智能的基础。此外脑科学的研究也为因果建模的突破提供了生物学的蓝图。例如中国科学院神经所王立平研究组发现猕猴通过整合先验信息和感觉输入推断隐藏的因果结构同时动态更新存储的先验知识和感觉表征来适应环境变化。帕金森病中副束旁核 Pf 解偶联的研究通过动态因果建模 DCM 量化 Pf 与其他脑区的有效连接为推动精准治疗策略的开发提供了因果框架。南方科技大学则通过扰动 AI 孪生脑首次刻画全脑有效连接图谱揭示因果连接的方向与强度。未来随着脑科学的发展人工智能的因果推断能力也将不断增强推动人工智能从“数据拟合”向“数据理解”演进。

矛盾发现能力揭示与调和认知冲突。矛盾性源于理论与观察或不同理论之间的根本性不一致。例如黑体辐射引发的“紫外灾难”作为理论预测与实验数据间的直接矛盾暴露了经典物理学的认知边界和内在局限。基于逻辑推理、跨域关联和动态演化人工智能独特

的“超逻辑”特性使其既能作为“矛盾揭示者”发现人类难以察觉的逻辑不一致或预测偏差主动暴露传统科学的认知盲区又能作为“矛盾解决者”提出新假设与调和的创新框架。以 Google DeepMind 的 GNoME 为例这一深度学习工具通过图神经网络预测了约 220 万种新晶体结构其中约 38 万种具备稳定性适合实验合成。该成果显著加速了材料发现产出相当于传统研究方法数百年的积累极大缓解了材料科学领域精度与效率难以兼顾的矛盾。传统量子化学中电子位置庞大的概率空间使得分子激发态无法精确计算而 FermiNet 神经网络突破这一局限以 4meV 误差解析碳二聚体结构将微观世界“概率云”与宏观可测性间的矛盾进行量化调和实现量子世界“不确定性”的具象化。未来人工智能作为认知进步的“催化剂”将通过整合高维复杂数据不断识别与调和内在逻辑冲突、修正和扩展人类认知体系持续推动跨学科的创新突破。

收敛逼近能力趋向复杂问题最优解。在复杂系统科学的指导下人工智能可以通过渐进式算法迭代寻求多维目标约束下的帕累托最优为实现系统性智能优化提供有力保障。强化学习中的 Q—Learning 算法通过泛函分析的 Banach 不动点定理证明了值迭代算法的收敛性确保在马尔可夫决策过程中逐步逼近最优策略。梯度下降法及其改进策略如动量法、自适应学习率则通过调整参数更新方向和学习率避免震荡或陷入局部最优从而提升收敛效率。人工智能算法的收敛逼近能力为工程系统中稳定性、泛化性和高效性的实现奠定有力基础。例如加文医学研究所研发的 AAnet 通过精准解析肿瘤细胞的多样化生物学模式实现从“一团乱麻”到肿瘤内单个细胞的清晰分型推动个性化联合疗法设计改进。药物研发领域原本需 2 ~ 3 年的抗艾滋病病毒 HIV 小分子设计的先导化合物筛选在人工智能驱动的科研平台辅助下科研人员 2 分钟就可生成超 25 万个全新分子 30 分钟便可筛选出 172 个潜在有效分子 25 实现了从“大海捞针”到“精准定位”的范式转变。

突变涌现能力带来“灵感式”创新。在特定情境下人工智能将

突破常规逻辑框架 产生超越训练数据范畴的原创性解决方案的能力。这种“灵光一现”并非随机偶然而是在算法机制、认知模型与复杂系统动力学的共同作用下由“量变”到“质变”的突破式涌现即通过复杂系统内部的结构重组自发出现超越组成部分功能的全新能力。这表明人工智能不再是仅停留于预设规则的机械执行而是实现了复杂系统中的自适应创新。2016 年 AlphaGo 在与李世石围棋对决的布局阶段放弃“外四路行棋”的千年定式将棋子置于第五路这一策略并非人类传授而是深度强化学习在博弈空间中的自发“突变式发现”奠定了 AlphaGo 的胜局。2023 年 Google Deep—Mind 的 FunSearch 框架面对经典组合数学问题中的高维“直觉依赖”瓶颈实现帽子集 Cap set 难题近 20 年来的最大规模突破。麻省理工学院推出的自主科学发现模型 SPARKS 则通过进行自主“假设—验证—改进”循环揭示了两条此前未知的蛋白质设计规则 长度相关的机械交叉解决了蛋白质设计折叠材料难以定量衡量的模糊性蛋白质设计的稳定性挫折区提供了一种避免脆弱的结构或改造失效的新手段。

综上关联识别能力到因果推理能力的实现是从“是什么”到“为什么”的转变增强了决策的可靠性矛盾发现能力与收敛逼近能力的结合通过从复杂约束中寻优实现了可持续的动态平衡突变涌现能力的颠覆式创新则催生出解决问题的新手段。五大能力并非孤立存在而是交织形成解决复杂问题的“智能网络”从而实现各领域赋能的提质增效。

2.4 系统整合

系统整合是人工智能实现智能系统与物理世界深度融合的工程化过程依赖于不同层级的技术实现。按照功能的划分人工智能技术涵盖包括基本理论方法与计算工具的基础技术、面向具体功能实现的功能技术以及针对特定行业或领域应用的领域技术。三类技术的有机整合实现人工智能从单点技术到高阶功能的跃升再到系统的工程化部署最终形成从“智能生成”到“智能执行”的闭环能力体系赋能经济社会各领域的智能化转型。

基础技术是人工智能的技术底座主要包括人工智能基础算法如机器学习、深度学习、强化学习等计算性能相关算法如高性能计算、分布式系统、联邦学习、隐私计算等以及人工智能学习框架等。基础算法作为人工智能技术的核心三大流派从单点突破到融合创新的演进催生出自适应的 AI 系统推动智能范式从感知向认知跃迁。计算性能相关算法则解决了 AI 的“效率瓶颈”高性能计算技术通过硬件与软件的协同优化让海量数据的并行处理成为可能。例如 DeepSeek—V3 通过 MLA 压缩内存、MoE 稀疏计算、FP8 精度优化及通信 — 计算重叠大幅压缩算力需求实现了高性能与低成本的平衡。人工智能学习框架作为连接算法理论与工程实践的桥梁通过封装底层算法逻辑与硬件接口让开发者无需深入理解矩阵运算、并行编程等细节就能实现模型的快速搭建。例如开发者使用 PyTorch 仅需 30 行代码即可完成 CNN 图像分类而同等功能的底层 CUDA 实现需上千行代码。这些技术为模型训练、推理执行、数据管理等环节提供关键支持支撑人工智能系统的可扩展性与模块化设计能力。例如作为 NLP 和多模态任务的基石 Transformers 模型通过自注意力机制高效处理序列数据支撑 ChatGPT 等大型语言模型 LLM 的对话、翻译和文本生成功能。基础技术的突破是实现高阶功能的前提也是推动人工智能从实验室走向工程化应用的关键条件。

功能技术是 AI 系统“可感、可认、可决”的关键包括计算机视觉、自然语言处理 NLP、语音识别与合成、情感计算等。功能技术的实现如同赋予机器“感官”与“大脑”使其能够从环境中获取信息、理解含义并做出反应。这些技术并非孤立存在而是构成了 AI“感知—推理—生成—决策”的完整功能链条在多模态交互、复杂任务处理等场景中通过协同配合让机器从被动执行指令升级为主动解决问题。例如自然语言处理 NLP 赋予 AI“语言理解与表达”能力搭建人机交互的桥梁 语音识别实现声音与文字的双向转换让人机交流更加自然流畅。基于此与 DeepSeek 深度融合的华佗 GPT 导诊系统通过分析患者症状描述实现精准科室推荐并一键挂号准确率达 90% 以上。

计算机视觉作为让 AI “看懂世界”的关键支持图像识别与目标跟踪。DeepDR—Transformer 辅助下的糖尿病视网膜病变 DR 转诊识别将使缺乏医疗资源和培训经验的初级保健医生 PCPs 的准确率由 81.0% 提升至 92.3%。AI 作为赋能人类的工具不仅通过“听懂”人类沟通信号、“理解”不同语言沟通需求、“看懂”显微镜下的病变细胞带来效率的提升更通过创新复杂问题解决方案拓展了人类能力的边界。随着多模态融合技术的成熟 AI 系统将实现更自然、更智能的人机交互从“帮人做事”迈向“与人协作”的新阶段。

领域技术是人工智能在具体行业中的工程化实现它如同连接 AI 能力与行业场景的桥梁将算法模型、数据资源与行业知识、场景逻辑、操作流程深度编织最终形成可落地、可复用、可迭代的解决方案。与基础技术的通用性和功能技术的工具属性不同领域技术更强调“场景适配性”——它不仅要追求算法在特定任务中的精度更需满足行业对部署便捷性、系统稳定性、结果可解释性的刚性需求。人工智能通过领域技术不仅实现技术与行业的深度融合还将形成人机协同、跨界融合的新形态通过人与 AI 的协同优化以及跨学科、跨行业的知识整合推动行业场景的智能化再造。这种“技术下沉”的过程正是人工智能从实验室成果转化为产业生产力的关键环节。以自动驾驶技术为例自动驾驶不仅依赖感知与决策等通用功能技术还需要与汽车工程、车辆控制系统、交通法规和城市路况数据深度耦合构建端到端的闭环自动控制系统。在这一过程中 AI 与驾驶员协同优化决策融合交通工程与城市规划知识形成跨界融合的智能驾驶生态。在医疗健康领域以影像识别、自然语言处理等功能技术为核心的电子病历、临床路径和医学影像集成将实现辅助诊断与个性化治疗。AI 与医生的协同诊断结合医学知识与患者数据的跨界整合将显著提升诊疗精准性与效率。从产业演进视角看领域技术的核心在于实现“技术—业务—治理”的三重整合领域技术的成熟度直接决定了行业数智化转型的深度。只有当 AI 真正理解行业的“规则”“痛点”与“红线”才能成为推动行业数智化转型的“主引擎”。这种深度融入行

业机理的技术形态将使人工智能从赋能提效的工具助手转变成支撑产业升级的“基础设施”。

总体来看未来随着人工智能技术的进一步发展人工智能赋能的数据整合—知识整合—系统整合将愈发紧密如催生代理型人工智能 Agentic AI 等实现“感知—推理—执行—学习”的闭环推动人工智能从计算智能向感知智能、认知智能、自主智能演进实现从自动化工具向自主系统的演进。

3. 人工智能赋能整合面临的问题挑战

2025 年 4 月 25 日中共中央政治局就加强人工智能发展和监管的第二十次集体学习中习近平总书记强调“面对新一代人工智能技术快速演进的新形势要充分发挥新型举国体制优势坚持自立自强突出应用导向推动我国人工智能朝着有益、安全、公平方向健康有序发展”。随着人工智能赋能应用向纵深发展人工智能在提升整体智能水平与应用效率的同时也带来了新的系统性风险和伦理困境等治理挑战。

3.1 数据整合面临机制缺陷与安全挑战

一是数据流通机制尚不健全。随着数据成为人工智能时代的关键性战略资源各方对数据的控制与利用日益重视出于安全保密、竞争优势等多重考虑政府部门、行业机构和企业普遍存在“重拥有、轻共享”的倾向导致跨部门、跨领域乃至跨区域的数据难以有效流通与整合。此外当前我国在数据确权、分类分级、合规使用、交易流通等方面的基础制度仍在逐步构建中缺乏统一的法律规范与激励机制使得大量数据资源处于“沉睡”状态真正能够被高效利用的数据占比偏低极大制约了人工智能系统的数据基础能力建设。特别是数据确权难成为阻碍数据整合的核心瓶颈尽管“数据二十条”创新性地提出了数据产权“三权分置”框架但在落地实施过程中仍缺乏操作细则和统一标准——如具体的权利内容、边界和行使条件尚不明晰各地在登记要求、主管单位等方面存在差异尚未形成互联互通

的跨区域统一数据管理平台等。二是数据标准存在鸿沟。长期以来不同部门和行业基于自身业务逻辑建立起独立的数据采集、分类、标注与存储体系形成各自为政的“信息孤岛”缺乏统一的接口与对齐机制不利于数据有效融合。尤其是在面对多源异构、结构非统一、语义不一致的数据场景时缺乏统一的语义规范与数据接口标准严重阻碍了人工智能系统对数据的理解、学习与调用能力。三是数据动态安全风险上升。人工智能赋能的数据整合使原本处于孤岛的数据得以连通形成“数据链”更编织成为“数据网”。然而随着数据在不同系统之间的频繁传输与共享数据静态安全风险向动态安全风险转变攻击者可能利用跨系统的漏洞进行更为精准的对抗性攻击带来个人隐私、关键制造工艺参数等敏感信息泄漏风险显著提升。

3.2 知识整合面临算法偏见和“黑箱”问题

一是算法偏见与决策偏差。模型训练高度依赖训练数据但数据本身隐含对不同群体、不同文化的差异如果训练数据集在不同对象间的分布失衡将进一步放大系统内嵌的数据偏差和模型误差导致技术中立难以实现。例如智能医疗助手使用中非标准口音人群因存在语音识别障碍而被边缘化女性和少数族裔则因代表样本不足而导致诊断准确率的降低。二是算法黑箱与不可解释性。人工智能赋能整合中复杂系统输入输出的深层次结构使得理解技术实施的内在路径变得愈加困难更使得“错误”源头难以追溯。这种“盲盒”则反过来进一步加剧了复杂系统的脆弱性。自动驾驶、智能安防、应急管理等领域隐秘的逻辑异常未被及时修正可能带来严重的技术失控风险造成无法挽回的后果。

3.3 系统整合面临安全脆弱性与伦理挑战

一是系统性安全风险增大。人工智能的工程整合导致系统高度复杂模块间依赖性增强风险攻击面扩大。一处故障或单点入侵可能通过整合链传导至整个系统诱发“级联失效”特别是模型、算力平台、传感器、通信系统等深度耦合一旦整合接口被攻击黑客可利用其通达多

个系统环节造成放大性破坏。二是概念验证到生产部署的鸿沟。尽管人工智能在实验室或封闭测试环境中往往展现出强大能力但一旦进入现实世界应用场景系统整合面临的数据分布偏移、边界条件不确定以及新旧系统间协调不足等问题带来原型落地难的问题。三是价值对齐难题导致伦理追责困境。人工智能系统整合涉及算法决策逻辑、数据价值偏好与人类伦理判断的深度耦合一旦缺乏统一价值指引将产生严重的伦理风险与责任归属困境。例如面对经典“电车难题”自动驾驶汽车无论做出何种选择都将面临道德的批判而技术的黑箱属性也导致责任归属存疑 更加剧事故后责任主体认定的困难。又如军事化应用潜藏巨大威胁人工智能武器远非完全智能无法进行精准的价值衡量一旦被滥用于极端作战方法和手段将带来毁灭性灾难。

4. 对策建议

4.1 完善数据流通制度和标准建设强化全生命周期数据安全防护

一是完善数据流通制度与激励机制。明确数据的所有权、使用权、交易权确保数据权利主体的合法权益为数据的流通和交易奠定基础。支持有条件的地区和重点行业开展数据确权授权区域试点建立容错与合规减责激励措施激发市场活力探索适合实际情况的数据授权模式打造一批示范样板。引导人工智能系统开发者采用隐私计算、联邦学习等技术确保模型训练过程中的数据“可用不可见”提高数据可使用性与安全性。推广区块链等可信技术应用利用其去中心化、不可篡改的特性实现数据资产唯一性确权与信息可追溯。根据数据的来源和贡献建立多方共赢的收益分配机制确保各方获得合理的回报激励数据提供者和处理者持续优化和利用数据。二是构建跨领域的的数据标准体系。制定推广数据采集、数据质量、目录分类管理、共享交换接口、共享交换服务等方面的标准确保数据的一致性和互操作性提高数据的可用性和共享效率。倡导行业协同治理设立工业互联网、智慧城市等“行业数据联盟”推动龙头企业开放核心数据接口标准形成标准统一、共享顺畅的行业数据生态。推动科研领域数据标准化与数据共享尤其是基因组学、气候科学、物理实验等领

域的数据促进跨领域数据能通过标准化接口高效融合使用促进跨机构、跨学科的数据共享和利用助力跨学科合作与创新。三是强化全生命周期数据安全防护。采用技术嵌入监管方式推广零信任架构与隐私计算技术对金融、医疗、科研等涉及个人信息的 AI 系统要求实施隐私保护设计认证。建立风险熔断机制对关键基础设施实施数据流动实时监控部署对抗攻击检测系统建立“数据安全熔断”预案一旦泄露自动隔离受影响模块。

4.2 加强算法偏差治理提升模型透明度与可解释性

一是注重推动 AI 基础研究。面向 AI 基础理论加强数学、物理学、统计学、心理学等多学科合作从源头加快 AI 新理论和新方法的创新。二是以数据和技术应对算法偏差。通过规范数据集对医疗、金融等高风险领域实施数据多样性审计。采取动态纠偏机制建立国家级算法测试平台提供偏见检测工具包及时纠正存在的算法偏差。三是破解算法黑箱提升透明性。针对深度学习的“黑箱”特性采用可视化分析、因果推理与规则提示相结合的多模态解释方式提升模型可解释性。推动 LIME、SHAP 等可解释人工智能 Explainable Artificial Intelligence XAI 技术在关键应用场景的深入研究与落地应用建立模型透明度评估标准与认证体系。引导科研人员负责任地使用 AI 鼓励在研究中使用 AI 模型时关注可解释性、透明性、可靠性等问题避免因算法不可解释带来对科研可信度的质疑。

4.3 创新智能系统整合落地提升工程韧性与伦理合规能力

一是弥合技术落地鸿沟。建立人工智能系统从概念验证到生产落地的中间验证平台提升模型在真实环境中的适应性。通过容器化部署与微服务架构提供可插拔算法组件与轻量级部署工具包支持在线调优与本地化定制降低基层机构在部署和运维中的技术门槛。二是构建韧性系统架构。制定《AI 系统整合安全规范》要求关键系统采用微服务架构通过“防火墙沙箱”隔离核心模块阻断级联失效。在智慧城市、自动驾驶等领域推广故障注入测试模拟硬件损毁、数据污染等极端场景强制年度压力测试。三是强化伦理治理与责任框

架。制定统一的人工智能伦理指南与审查框架细化数据采集、算法设计、结果应用等各环节的责任边界与合规底线。组建跨学科、多方参与的伦理审查委员会融合法律、社会学、计算机科学等专业视角对高风险项目实施常态化审查引入可量化的风险评估指标与分级管理制度实现对潜在伦理隐患的事前预警与事中监控。在科研、医疗等领域鼓励相关科研院所建立伦理委员会或伦理审查制度确保 AI 应用中的伦理标准得以遵循确保科研决策符合科学规范、伦理标准和社会价值观。

本文从人工智能的发展演进出发阐释了人工智能赋能的内在机制即“数据—知识—系统”三重整合。但需要认识到在人工智能赋能的基础上还需要关注“效率”和“安全”的问题如人工智能技术应用是否带来单点效率与系统效率悖论以及人工智能赋能各领域带来的衍生安全风险等问题。为此未来还需要从管理学视角出发进一步结合人工智能赋能各领域的实际应用场景 深入探讨人工智能赋能过程中的组织适配与治理机制构建问题以实现效率与安全动态平衡的人工智能赋能闭环。

作者信息

潘教峰 第十四届全国人大代表、中国科学院大学公共政策与管理学院院长、中国科学院科技战略咨询研究院院长

王楚扬 中国科学院大学公共政策与管理学院硕士研究生

吴 静 中国科学院科技战略咨询研究院研究员

周洪宇

技术、制度与思想：生成式人工智能在教育领域中应用的演进逻辑

来源 | 《电化教育研究》2024 年 08 期



第十三届全国人大常委会委员、中国教育学会副会长、长江教育研究院院长兼华中师范大学国家教育治理研究院院长 周洪宇

一部教育史就是一部教育与技术交融的互动史。一种具有里程碑意义的新技术的出现，不仅能够更新人类生活方式的面貌，而且可以促进人类教育方式的变革。历次教育变革均为新技术在教育领域中广泛应用的结果，教育变革反过来又进一步推动了教育新技术的发明与创新。新技术在教育领域中的应用，往往是先有其事，后有其名，再有其学；常常是先有技术，后有制度，再有思想。教育技术、教育制度与教育思想是教育文化的三项基本要素，不同要素在教育文化中具有不同的地位和作用，它们的关联构成了教育文化的基本结构。审视当下的教育生态，以 ChatGPT 为代表的生成式人工智能正在介入教育领域，正在引发教育的变革。为了更好地应对生成式人工智能在教育领域中应用的机遇和挑战，我们必须从技术、制度、思想三个层面出发，把握好生成式人工智能在教育领域中应用的形态转换和演进逻辑。

二、技术层面的冲击与回应：生成式人工智能在教育领域中的应用

技术是物质要素，既是发展的对象，也是发展的结果，它制约发展并影响着制度和思想。人类在面临科技革命的浪潮时，总是以技术优化教育、以教育改进社会作为回应。毫无疑问，生成式人工智能在教育领域中具备巨大的应用潜力，但从生成式人工智能当前展现的功能来看，尚无法满足正常的教育教学场景的需要。生成式人工智能更多的是让我们洞见未来教育的可能性，一种作为争议性的、变革性的、辅助性的、缓慢性的教育技术的可能。

（一）基于算法决策而产生的矛盾：争议性的教育技术

教育问题是教育的“特色”，每当新技术闯入教育实践活动，都有数不尽的新问题产生。回顾教育技术的发展史，每当新技术被应用到教育领域时，都会在教育界引起广泛的热议和普遍的担忧。当电视机、投影仪、摄像头、计算机、平板电脑、智能手机进入校园时，同样引起了一系列有关教育公平、身心健康、师生关系、数据隐私、技术依赖等问题。ChatGPT 爆火后，生成式人工智能的应用前景在我国教育界引发了大讨论。

生成式人工智能模型集强算法、大算力、大数据于一体，是生成式人工智能的关键组成部分。生成式人工智能的崛起是由算法、算力、数据三个方面共同进步所推动的。其中最重要的是算法方面的进步，而个体对算法的态度，尤其是对新生算法的态度总是具有矛盾的两面性：算法欣赏、算法厌恶，甚至欣赏与厌恶并存。算法欣赏与算法厌恶代表个体对于人工智能算法的态度，它们是基于算法决策而产生的。

持算法欣赏态度的教育工作者倾向于借助算法来做出决策，他们相信生成式人工智能可以为教育领域带来积极的影响和改进，从而推动教育的数字化、自动化、普及化和均衡化。持算法厌恶态度的教育工作者则对算法或算法应用感到反感或持消极态度。他们认为生成式人工智能的设计或实现存在一些争议，生成式人工智能进入教育领域可能会导致教育异化、加剧数字贫困、侵犯知识产权、陷入数据偏见、传播虚假知识。教育工作者围绕生成式人工智能而表现出来的欣赏或厌恶，是一种基于算法决策而产生的矛盾心理，这种矛盾心理促使生成式人工智能教育应用成为极富争议性的教育话题。

（二）基于劳动解放而带来的助力：变革性的教育技术

劳动解放是人的全面发展的内在尺度，是马克思人类解放学说的核心内容。马克思认为，人是能够制造与使用工具以从事生产劳动的动物。人类制造与使用工具的目的是弥补自身劳动能力的不足，减轻自身的劳动负担。从古至今，无论是石器时代、青铜时代、铁器时代、蒸汽时代、电气时代、信息时代，抑或是人工智能时代，人类都希望通过制造和使用工具，在提高生产效率的同时，不断解放自身的体力劳动和脑力劳动。从工具的发展历程来看，生成式人工智能也是人类制造与使用的一种新型工具。生成式人工智能将对人类的脑力劳动解放产生重要的推动作用，也将为人的全面发展提供必要的技术准备和物质基础。

作为一种变革性的教育技术，生成式人工智能对教育的底层逻辑产生了极大的冲击，尤其是在知识生产和人才培养两个方面。从知识生产来看，生成式人工智能旨在利用语言模型，模仿训练数据内容来生成相似的知识。学习者和研究者习惯于通过生成式人工智能来完成家庭作业和撰写学术论文。例如：美国的一项调查显示，超过 85% 的受访学生会使用 ChatGPT 来完成家庭作业；安东尼·奥曼教授则发现，全班排名第一的论文竟然是用 ChatGPT 写的。虽然生成式人工智能对知识生产的冲击比人们预想中更大，但是生成式人工智能使用的是一种复制常见模式的预测算法，不具备原创性且缺乏风格，生成的只是计算性的、重复性的、机械性的、非创造性的知识，无法生成有情感的、独特的、灵活的、创造性的知识。因此，生成式人工智能只是为人类的脑力劳动解放提供了有限度的选择，并不能完全实现人类脑力劳动的解放。从人才培养来看，生成式人工智能有效地推动了人才培养标准的转向，传统的知识型人才逐渐不能满足时代的需要，复合型人才和创新型人才正在成为人才培养的主流。

（三）基于技术先行而衍生的实验：辅助性的教育技术

随着生成式人工智能的快速发展，世界各地的教师和机构都开始尝试开展小范围的生成式人工智能辅助教育实验。《麻省理工科技评论》AI 资深编辑威尔·道格拉斯·海文（Will Douglas Heaven）访谈了一些将 ChatGPT 应用到教学实践中的美国教师，他们认为经过几个月的课堂实验，生成式人工智能不仅是作弊者的梦想机器，实际上也可以作为强大的课堂辅助工具。从焦虑、恐惧到部分接受，再

到谨慎实验，教育工作者在同生成式人工智能接触的过程中，不断地挖掘生成式人工智能的应用潜能，利用生成式人工智能辅助教育事业发展成为世界各国的普遍共识。

作为一种辅助性的教育技术，生成式人工智能的辅助功能涉及学习、教学、科研等方面。第一，生成式人工智能可以帮助学生搭建自主学习平台，为学生量身定制学习计划和学习内容，自动化评估学生的学习状况，充分满足学生个性化成长的需求。例如：Quizlet 依托 ChatGPT，增加了一个名为 Q—Chat 的功能，可以根据学生的学习偏好定制学习内容，为学生提供类似于一对一的定制教育服务。学生既可以选择不同版本的学习材料，也可以及时调整学习材料的难易程度。第二，生成式人工智能可以辅助教师开展教学工作和事务性工作，提高教师在课程计划制定、教学大纲设计、作业批改、考试评分、活动策划、汇报撰写、报表填写等方面的工作效率，全方位地减轻教师的压力和负担。第三，生成式人工智能具备强大的科研辅助能力，可以帮助学生撰写与修改论文。一方面，生成式人工智能可以根据相关主题生成文献综述，帮助学生快速了解相关领域的研究进展；另一方面，生成式人工智能可以修改完善论文的句式、语法和结构，帮助学生提升论文的整体质量。

（四）基于系统融合而出现的挑战：缓慢性的教育技术

当今世界正经历百年未有之大变局，新一轮科技革命与产业变革深入发展，科技进步与教育发展的融合态势已初露端倪。然而，教育系统与科技系统的融合存在高度的差异性和不确定性。从电化教育在中国传入与发展的近百年历程来看，技术促进教育系统的变革过程往往是复杂且缓慢的。教育深层变革的真正实现需要走过漫长的系统融合历程，需要解决一系列的技术集成问题、资源分配问题以及教育者与学习者的适应性问题。

作为一种缓慢性的教育技术，生成式人工智能在引入教育领域的过程中遭遇了诸多挑战。一方面，挑战来自教育改革的内生动力不足。大部分教育工作者面对新技术习惯于采取保守的态度，几乎不愿意走出舒适区，这就导致教育变革的速度往往滞后于技术更新的速度。“教育—社会—科技”的“内驱轮”似乎与“科技—社会—教育”的“外驱轮”的耦合机制不协调。因此，当生成式人工智能展现出

教育应用的潜能后，大部分教育工作者首先感受到的却是恐惧和担忧。另一方面，挑战来自生成式人工智能自身存在的局限性。一是生成式人工智能生成的内容错误率较高，存在误导学生的风险；二是生成式人工智能生成的内容有时存在偏见，可能会对教育活动产生负面的影响；三是生成式人工智能生成的内容缺乏对学生批判性思维和创造性思维的养成，可能会导致学生过度依赖生成式人工智能生成的内容。

三、制度层面的规范与引导：生成式人工智能在教育领域中的发展

制度是组织要素，是影响技术发展的重要因素，约束、激励、调节着技术的发展方向。制度的运行由不同的层级构成，美国著名政治学家埃莉诺·奥斯特罗姆（Elinor Ostrom）将制度运行的层级分为法律层、管理层和操作层。从埃莉诺·奥斯特罗姆的“制度层级分析理论”来看，生成式人工智能在教育领域中的发展，离不开以上三种层级的逐层衔接，只有形成耦合完整的生成式人工智能教育制度，才能产生最好的规范与引导作用。

（一）立法保障：生成式人工智能教育的法治建设

生成式人工智能教育制度的法律层是指中央政府或地方政府开展生成式人工智能教育治理的法制性规定，具有最高层次的权威性和约束力。此层级是决定生成式人工智能在教育领域中应用发展的元规则，赋予生成式人工智能以合法性，为生成式人工智能营造良好的制度环境，使生成式人工智能在不违背法律法规的前提下运行，解决生成式人工智能宏观保障缺失问题。

1. 关于生成式人工智能的法治建设

随着生成式人工智能的快速发展，各国政府一直围绕生成式人工智能应用开展法治建设，旨在有效监管生成式人工智能在教育等社会领域的应用，解决生成式人工智能带来的诸多科技伦理问题。近年来，包括中国、美国、欧盟在内的国际治理主体，都在快速推进关于生成式人工智能的法治建设。2022年，美国先后颁布了《2022年算法问责法案》《美国数据隐私和保护法案》《人工智能权利法案蓝图》，在现有的立法框架内对数据和算法进行监管，旨在保护数据隐私，防止算法歧视，促进社会公平。2023年7月，国家网信办等七部门联合发布了《生成式人工智能服务管理暂行办法》，旨在促进生成式人工智能的健康发展和规范

应用，维护国家安全和社会公共利益，保护公民、法人和其他组织的合法权益，凸显了中国特色的生成式人工智能治理之道，为全球生成式人工智能治理贡献中国智慧。2023 年 12 月 8 日，欧洲议会、欧盟成员国和欧盟委员会三方就《人工智能法案》（AI Act）达成协议，该法案将成为全球首部人工智能领域的全面监管法规，将为全球范围内有效监管人工智能技术定下基调，标志着全球人工智能立法进程迈出了重要一步。

2. 关于生成式人工智能的法律法规

各国关于生成式人工智能的法律法规对生成式人工智能在教育领域中的应用提供了相关的法律参照。例如：ChatGPT 存在使用受版权保护的材料来训练人工智能模型的现象，学生在使用 ChatGPT 辅助论文写作时，常常会遭遇侵犯他人知识产权的问题。中国和欧盟为了应对生成式人工智能带来的版权挑战，规范生成式人工智能所输出的内容，在版权保护方面都实施了监管举措。《生成式人工智能服务管理暂行办法》要求生成式人工智能供应商对生成式人工智能模型训练过程中使用的受版权保护的内容进行数据标注。《人工智能法案》则要求生成式人工智能开发人员披露他们在构建系统时使用的受版权保护的材料。关于生成式人工智能教育的法治建设，一方面应依托宏观的生成式人工智能法律法规，另一方面应制定专门的生成式人工智能教育法律法规，全方位地提供立法保障。

（二）政策管理：生成式人工智能教育的行政规范

生成式人工智能教育制度的管理层是指通过政策的制定和评价使法律层进一步制度化，进而为操作层提供通用性规则，以保证制度层级的上下贯通。此层级是生成式人工智能在教育领域中应用发展的通用性规则，由政府行政机关制定的各种规章、规定构成，保证生成式人工智能在政府的管理和监督下运行，解决生成式人工智能政府管理缺失问题。从各国政府对人工智能的治理经验来看，可以从三个方面出发来规范生成式人工智能，提升生成式人工智能教育治理的现代化水平。

1. 制定国家人工智能发展规划

制定国家人工智能发展规划，实施国家人工智能发展战略，牢牢把握人工智能发展新阶段国际竞争的战略主动性。联合国教科文组织指出，截至 2023 年初，

全球大约有 67 个国家制定了关于人工智能的国家规划。早在 2017 年，我国便发布了《新一代人工智能发展规划》，强调完善人工智能教育体系，建设人工智能学科，形成人工智能教育高地和人才高地。根据国家人工智能发展规划，教育行政部门应准确把握生成式人工智能发展态势，找准突破口和主攻方向，拓展生成式人工智能应用深度与广度，全面提升教育智能化水平。

2. 制定国家通用数据保护条例

制定国家通用数据保护条例，规范数据使用者对国家通用数据的收集、传输、保留和处理，明确数据使用者的责任和义务，保护数据主体的权利和隐私，避免数据使用者代替数据主体成为实际的数据控制者。根据国家通用数据保护条例，教育行政部门应规范生成式人工智能供应商对教育数据的收集和处理，要求生成式人工智能供应商依法开展预训练、优化训练等训练数据处理活动，防止教育数据的泄漏与滥用，维护教育数据的安全与稳定。

3. 制定人工智能科技伦理治理办法

制定人工智能科技伦理治理办法，完善人工智能科技伦理治理体系，有效防控人工智能科技伦理风险，提升人工智能科技伦理治理能力，努力实现人工智能技术高质量发展和高水平安全良性互动。2022 年 3 月，中共中央办公厅、国务院办公厅发布了《关于加强科技伦理治理的意见》，提出了我国科技伦理治理的原则以及基本要求，不断推动科技向善、造福人类，实现高水平科技自立自强。根据人工智能科技伦理治理办法，教育行政部门应完善生成式人工智能教育伦理规范和标准，建立生成式人工智能教育伦理审查和监管制度；落实生成式人工智能供应商科技伦理管理主体责任，引导供应商和研发者自觉遵守生成式人工智能教育伦理要求；加强生成式人工智能教育伦理理论研究，提升生成式人工智能教育治理水平。

（三）操作指引：生成式人工智能教育的行动指南

生成式人工智能教育制度的操作层是指通过具体的实践行动使管理层进一步细化，为生成式人工智能在教育领域的真正落地提供具体的操作细则，使生成式人工智能教育制度层级体系具备可行性基础。此层级是生成式人工智能在教育领

域中应用发展的基础性规则，规定了生成式人工智能教育如何改善课程、教学、学习的行动指南，解决生成式人工智能教育微观执行缺失问题。生成式人工智能在教育中的作用需要在实践中得以展示，可以从三个方面出发来提升生成式人工智能教育的实践活力。

1. 制定适应生成式人工智能的课程改革方案

制定适应生成式人工智能的课程改革方案，打造智能、融通、开放的课程模式，加强课程与生成式人工智能的连接，更好地迎接生成式人工智能对课程建设提出的挑战。在课程目标设置上，应聚焦学生智能素养培育，重点关注人的智能发展以及运用智能工具解决问题的能力。在课程内容设置上，应聚焦学生复合素养培育，倡导知识和技能的跨学科整合，帮助学生掌握多元的知识和技能。在课程形态设置上，应聚焦学生创新素养培育，强调标准化课程与定制化课程相结合，努力实现“一人一张课程表”。

2. 制定适应生成式人工智能的教学改革方案

制定适应生成式人工智能的教学改革方案，打造人机协同、虚实结合的教学模式，不断改进和完善教学环境，以适应科技的发展和学生的需求。在教学目标设置上，应着重培养学生的思维意识和创新意识，提升学生掌握和应用生成式人工智能的能力，进一步巩固学生学习的主体性地位。在教学内容设置上，应将生成式人工智能的相关知识和技能纳入教学内容中，让学生充分了解生成式人工智能的基本原理、应用场景和发展趋势。在教学方式设置上，应采用人机协同的教学方式，利用生成式人工智能为学生提供更加个性化和高效的学习体验。在教学评价设置上，应建立基于生成式人工智能的教学评价体系，通过数据分析和智能评估，帮助学生改进学习方法，提高学习效率。

3. 制定适应生成式人工智能的学习改革方案

制定适应生成式人工智能的学习改革方案，打造自主性、独特性、多样性、创新性、终身性的学习模式，帮助学生树立内驱发展、个性发展、全面发展、创新发展、终身发展的学习观。在学习目标设置上，应围绕学生发展核心素养，重点培养学生的自主学习能力、自我认知能力、创新实践能力、数据分析能力和终

身学习能力。在学习内容设置上，应提供定制化、个性化、智能化的学习资源，不断激发学生的学习兴趣和学习欲望。在学习方法设置上，应鼓励学生使用问题式、项目式、协作式、探究式的学习方法，为学生的成长提供无限的可能。

四、思想层面的更新与重塑：生成式人工智能在教育领域中的普及

思想是观念要素，思想与技术的融合程度决定着技术在教育领域中的普及程度。技术与教育的深度融合必然蕴含着思想层面的更新与重塑，需要不断突破已有的观念限度，持续推动教育观念的发展与进步。生成式人工智能在教育领域中的普及必须紧紧依靠思想层面的更新与重塑，主要包括知识观、价值观、人才观的更新与重塑。

（一）树立开放、共享、自主的知识观

知识观是对知识的本质、价值、生产、传播等问题的基本认知，影响着人们的学习、思考和教育方式。面对生成式人工智能重塑知识观的现实，教育必须及时在思想层面上做出思考与回应，树立开放、共享、自主的知识观，以适应生成式人工智能发展和教育数字化转型。

1. 开放的知识观

开放的知识观强调多元性、动态性、丰富性，旨在帮助学生适应持续迭代的教育技术和不断变革的教育形态。树立开放的知识观，应帮助学生尊重与接纳不同领域、不同文化背景的知识，开阔心胸与视野，加深知识理解；应帮助学生打破思维定式，接受新知识，学习新知识，更新知识结构，提升适应能力；应帮助学生养成终身学习的意识，秉持持续学习的态度，丰富知识体系，夯实文化基础，逐渐成为全面发展的人。

2. 共享的知识观

共享的知识观强调公共性、交流性、合作性，旨在鼓励学生通过知识共享平台获取和分享知识，以知识的流通和增值促进社会的进步和发展。树立共享的知识观，应搭建公共知识共享平台，强化知识的公共资源属性，为学生共享知识提供公平、便捷的渠道；应帮助学生积极参与知识交流活动，通过知识的碰撞，让

学生的思维之光更加耀眼；应鼓励学生与其他知识主体协同合作，有组织地解决复杂的学习问题，实现知识主体的互利共赢。

3. 自主的知识观

自主的知识观强调主动性、独立性、原创性，旨在倡导学生在学习知识的过程中充分发挥自身的主观能动性，倡导学生在独立思考和积极探索的过程中形成具有原创性的知识体系。树立自主的知识观，应培养学生具备主动的学习能力，帮助学生主动寻找适应时代发展和自身需要的学习内容和方法；应培养学生具备独立的思考能力和批判能力，帮助学生理性地对待生成式人工智能所生成的知识，避免学生过度依赖并被单一的知识源所束缚；应培养学生的创新能力，帮助学生提出并验证新观点、新见解，充分激发学生的创造潜能。

（二）树立安全、公正、诚信的价值观

价值观是指人们对人与外界物的价值关系的认识，并在此基础上确定人类行为的价值取向。生成式人工智能与教育主体之间价值关系的特点包括多元性、变动性和可逆性。为了更好地应对生成式人工智能带来的机遇和挑战，应树立安全、公正、诚信的价值观，扬长避短，正确处理好生成式人工智能与教育主体之间的价值关系。

1. 安全的价值观

安全的价值观强调人本性、责任性、保护性，旨在通过安全至上的价值导向，促进生成式人工智能发展与教育主体生命成长的双和谐。生成式人工智能是新时代影响教育安全的显著变量，它一方面为教育领域带来巨大福利，另一方面也可能给教育安全带来极大风险，有效预防风险是确保安全的关键。因此，应将教育主体的生命安全作为生成式人工智能在教育领域中应用的第一原则，加强开发者、使用者和监管者的安全责任意识，加强隐私保护和数据保护，时刻关注并降低使用风险，确保安全成为生成式人工智能发展与治理的首要关切。

2. 公正的价值观

公正的价值观强调公平性、合理性、正义性，旨在防止生成式人工智能在教育领域的应用中产生歧视和偏见，确保生成式人工智能教育公平惠及所有学生。

树立公正的价值观，应对教育科技公司所提供的生成式人工智能算法进行严格的审查与测试，以公开透明消除生成式人工智能算法运行的灰色地带，确保生成式人工智能算法的公平性；应制定正当的教育政策，科学指导生成式人工智能教育资源分配；应鼓励多元主体共同参与生成式人工智能教育治理，秉持正义立场，维护公正的教育秩序。

3. 诚信的价值观

诚信的价值观强调真实性、准确性、可靠性，旨在塑造一个诚实守信的教育环境，为生成式人工智能教育的发展提供坚实的道德基础和社会支持。树立诚信的价值观，应鼓励教育主体在使用生成式人工智能的过程中提供真实的教育数据，以使生成式人工智能能够提供有效的评估、互动与指导；应要求技术开发者持续优化生成式人工智能，尽可能保证信息生成的准确性；应加强对生成式人工智能的法律监管、系统测试和模型训练，综合提升生成式人工智能的可靠性，为教育应用提供更加可靠的技术支持。

（三）树立新质、复合、智慧的人才观

人才观对教育观具有指导和引领作用，教育观是人才观的具体实践和体现，人才观和教育观在相互作用中不断完善和发展。在生成式人工智能时代，应树立新质的、复合的、智慧的人才观，以新质人才观塑造新质教育观，加快拔尖创新型人才培养，赋能新质生产力发展。

1. 新质的人才观

新质的人才观强调引领性、颠覆性、战略性，旨在培养具备引领思维、颠覆意识和战略视野的高素质人才。树立新质的人才观，应注重培养学生具备引领社会发展和行业进步的能力，培养学生具备颠覆传统、打破陈规的能力，培养学生具备战略规划和长远布局的能力，从而不断推动生成式人工智能技术的迭代升级，推动生成式人工智能教育的全面普及。

2. 复合的人才观

复合的人才观强调跨界性、全面性、创新性，旨在培养具备交叉素养、综合

素质和创新思维的高素质人才。树立复合的人才观，应注重培养学生具备跨领域、跨学科、跨行业的能力，培养学生具备文化基础、自主发展、社会参与的能力，培养学生具备敏锐观察、勇于开拓的能力，从而不断推动生成式人工智能与教育领域的知识、技术和资源进行有效整合，为解决复杂的生成式人工智能教育问题提供有效的思路 and 方案。

3. 智慧的人才观

智慧的人才观强调适应性、前瞻性和可持续性，培养具备快速适应变化、洞察未来趋势和可持续发展的高素质人才。树立智慧的人才观，应注重培养学生具备强大的学习能力、适应能力，不断吸收新知识、新技能，以应对不断变化的教育环境和社会环境；培养学生具备洞察趋势、预见未来的能力，把握生成式人工智能时代的发展方向；培养学生具备终身学习、不断成长的能力，以人才的持续发展推动生成式人工智能的持续进步。

文化的技术层是最活跃的因素，变动不居，日新月异；文化的制度层是最权威的因素，规定着文化整体的性质；文化的思想层是最保守的因素，是文化成为类型的灵魂。文化的变化，首先是技术层面的冲击与回应；习之既久，渐可认识制度层面的规范与引导；最后，方能体味思想层面的更新与重塑。生成式人工智能在教育领域中的应用是一个涉及技术、制度与思想多个层面复杂的、融合的、递进的过程，只有在这三个层面都实现深度融合和创新，才能充分发挥生成式人工智能在教育领域中应用的潜力，推动教育的数字化转型和升级。

作者信息

齐彦磊 河南大学教育学部副教授、博士

周洪宇 第十三届全国人大常委会委员、中国教育学会副会长、长江教育研究院院长、华中师大国家教育治理研究院院长

杨宗凯

推动人工智能融入教育全要素全过程全场景

来源 | 《中国教育报》2026年4月20日



教育部教育数字化专家咨询委员会主任、中国教育发展战略学会副会长 杨宗凯

人工智能正深刻改变知识生产、传播与应用的方式，也在持续重塑教育的组织逻辑、供给形态与治理模式。“人工智能+教育”的核心在于两个字，一是“融”，二是“变”。“融”，是立足于人工智能与教育深度融合格局基本形成的发展目标，让人工智能进入教育运行的基本结构、基本链条和基本空间；“变”，则是充分发挥人工智能赋能教育变革的引擎作用，推动办学模式、育人方式、管理体制和保障机制系统性转变。其中，全要素、全过程、全场景，是“融”的具体展开；办学、育人、管理与保障，构成了“变”的实践落点。

“全要素”，即推动人工智能深度融入教育关键要素。教育并不是某一个单独环节的简单叠加，而是由学生、教师、环境等多个方面共同构成的复杂系统。《“人工智能+教育”行动计划》（以下简称《行动计划》）在部署中并未停留于静态罗列，

而是突出人工智能对教、学、管、研、国际化等关键节点的赋能作用。如在教师方面，强调全面提升数字素养和智能应用能力，推动教师借助智能技术优化教学设计、改进教学方法、提升专业水平。在学生方面，强调加强人工智能通识教育，提升学生面向智能时代的学习能力、思维品质和复杂问题解决能力。在学科专业和科研方面，则要求顺应知识生产方式变革，推动学科交叉融合、专业动态调整和科研范式创新。因此，“全要素”所强调的，不是将技术简单叠加于既有教育结构，而是推动人工智能与教育关键要素发生更深层次的耦合，进而增强教育体系的适应性、韧性与创新能力。

“全过程”，即推动人工智能要贯穿教育发展的各个学段、各个阶段。《行动计划》部署了人工智能全学段教育和全社会通识教育：中小学阶段重在普及人工智能通识教育，夯实数字素养与认知基础，帮助学生建立对人工智能的基本理解和正确态度；职业教育阶段重在对接产业需求，推动专业升级和技能重构，提升学生面向智能生产、智能服务和智能管理的实践能力；高等教育阶段重在强化基础研究、交叉融合和拔尖创新人才培养，推动人工智能成为公共基础课和学科交叉的重要支撑；终身教育重在面向社会成员开展通识教育和技能培训，提升全民智能素养，增强适应技术变革的能力。由此，形成由低到高、由基础到专业、由在校到社会的贯通式发展格局。进一步看，学段贯通并不只是时间顺序上的延展，更是教育目标、课程内容、培养方式和评价机制的系统衔接。人工智能不仅要进入每一个学段，更要根据不同学段的教育功能和发展任务，形成循序渐进、螺旋上升的培养体系。

“全场景”，意味着人工智能教育应用必须突破传统课堂和学校边界，进入更加开放、复合、协同的教育空间。随着学习方式、资源形态和育人组织方式的深刻变化，教育活动日益发生在学校、家庭、社会、网络空间以及虚实融合环境的交互中。《行动计划》提出建设未来课堂、未来学校、未来学习中心和未来实训中心，推动主题式学习场景、虚拟仿真实验、智慧慕课和智能终端协同应用，其核心旨趣就在于构建多维联动的新型教育生态，使学习活动能够在更加真实、更加丰富、更加个性化的场景中展开。“全场景”不仅意味着技术应用空间的拓展，更意味着教育组织方式和资源供给方式的重构。尤其值得关注的是，文件对乡村学校、边远地区、特殊教育群体和社会学习者给予了专门部署，表明人工智能教育应用并非仅服

务于高水平学校和优势地区，而具有鲜明的普惠导向，即通过技术手段降低优质教育资源的获取门槛，促进教育公平从机会层面走向过程与质量层面。

从实践推进的角度看，推动人工智能融入教育全要素、全过程、全场景，实质上要求实现 4 个方面的深层转变。

其一，在办学模式上，实现从相对封闭、单一供给的传统办学，向开放共享、协同联动、无边界连接的智能办学转变。人工智能带来的并非某一教学环节的小修小补，而是学校组织形态和资源供给体系的深刻重构。过去，学校更多依赖校内课程、校内教师和校内空间开展办学；而在智能时代，优质教育资源将以更加广泛的形式流动起来，在更大范围内链接课程、教师、平台和产业资源。学校要从资源占有型办学转向资源整合型、平台支撑型办学，推动形成政府、学校、企业、科研机构和社会多元参与的资源供给格局。尤其是在职业教育和高等教育领域，应进一步推动科教融汇、产教融合，探索校企协同开发课程、共建项目、共育人才的新机制，让学校教育更加贴近科技前沿、产业变革和真实问题情境。可以说，人工智能正在推动办学模式从边界清晰的封闭系统，走向开放共享的生态系统，这本质上是办学组织的再造，也是教育改革的重要突破口。

其二，在育人方式上，实现从知识传授主导向个性化学习、能力导向和人机协同育人转变。智能时代，知识获取的门槛显著降低，单纯围绕知识点讲授和标准答案训练的育人方式，已难以适应未来对人才的要求。教育的重心必须从知识传授转向能力培养，从整齐划一的进度安排转向因材施教与个性发展。人工智能可以通过学情分析、路径推荐、过程性评价等方式，为学生提供更加精准的学习支持，也为教师开展差异化教学、实施精准育人创造条件。未来课堂应更加重视 PBL 等问题导向、项目导向的学习方式，推动“师—机—生”三元协同，促进学生在真实任务和复杂情境中学习、探究、合作与创造，更加注重培养学生的选择判断力、深度提问力、重构创新力。同时，人工智能可以分担重复性、程序性任务，却不能替代教师在价值引领、情感支持、伦理判断和人格塑造中的核心作用。因此，育人方式的变革，必须把未来教师和未来课堂建设紧密结合起来，推动教师从知识传授者更多转向学习设计者、成长促进者和创新引导者。

其三，在管理体制上，实现从层级式、经验式、相对粗放的传统管理，向扁平敏捷、数据驱动、精准决策的现代治理转变。长期以来，教育管理在不少场景中仍较多依赖经验判断、静态统计和分段管理，往往存在响应不够及时、协同不够顺畅、配置不够精准等现实挑战。人工智能的引入，为教育治理流程再造提供了重要条件。《行动计划》以教育智能大脑牵引人才供需、考试评价、就业服务和安全预警等领域改革，推动治理从碎片化走向整体化、从经验驱动走向数据驱动、从事后反应走向前瞻研判。未来学校不仅是技术设备更多、应用场景更丰富的学校，更应是治理结构更科学、管理运行更高效、决策机制更精准的学校。通过智能分析和辅助决策，学校和教育行政部门可以更好把握学生发展规律、优化资源配置、调整专业结构、提升管理效率，推动管理体制向更加扁平、敏捷、协同的方向演进。当然，人永远是治理的主体，要始终坚持技术辅助决策、责任最终归人，在提升治理能力的同时，守住教育公平、教育规律和伦理底线。

其四，在保障机制上，实现从分散建设、局部支撑、要素叠加，向制度统筹、标准牵引、平台贯通、协同创新的系统保障转变。人工智能深度融入教育，必须形成覆盖政策、标准、基础设施、师资培养、科研支撑、安全治理和生态协同的完整保障体系。一方面，要加强算力、数据、平台、模型等新型教育基础设施建设，健全数据治理、算法规范、隐私保护、内容安全和风险防控机制，为教育智能化提供可靠底座。另一方面，要把教师队伍建设摆在更加突出的位置，围绕人工智能赋予教师的新角色、新使命，必须提升教师驾驭智能工具、优化教学流程、开展人机协同育人的能力。同时，要加强高校、科研机构、政府、企业、学校之间的协同联动，形成多方支持的 UGBS 协同创新模式，在基础研究、技术研发、场景落地、评价改革等方面形成合力。只有通过制度保障、资源保障、师资保障和创新生态保障协同发力，才能推动“人工智能+教育”从试点探索走向规模应用、从技术可用走向教育好用。

邬志辉

生成式人工智能时代高等教育变革： 本体挑战、异化风险与体系重构

来源 | 《湘潭大学学报（哲学社会科学版）》2026年02期



东北师范大学中国农村教育发展研究院教授、博士生导师 邬志辉

教育、科技、人才是全面建设社会主义现代化国家的基础性、战略性支撑，高等教育作为三者的关键结合点，在推进中国式现代化进程中发挥着不可替代的龙头作用。当前，以生成式人工智能为代表的技术革命浪潮席卷全球，既为知识生产方式的变革带来效率曙光，也为传统高等教育体系带来了挑战。当生成式人工智能在知识整合与学术测验中展现出比肩甚至超越人类的能力时，传统高等教育培养的能够精准记忆、复述和应用知识的人才，正面临被算法替代的风险。这种外部技术冲击与内部教育困境的碰撞，迅速引发了全球范围内的激辩与审思，乐观者将其誉为开启个性化学习的“金钥匙”，认为它有潜力成为教育的强大辅助工具；悲观者则担忧其可能引发学术伦理失范、思维能力退化乃至教育公平异化等危机。与此同时，全球多国亦纷纷引入监管框架，以防范其潜在风险。这场

激辩反映了人类社会在面对颠覆性技术时的共同迷茫，也迫使我们回归对高等教育本源与价值根基的深层追问：在算法擅长“计算”的时代，人类“算计”的智慧又将何以安顿？当知识获取变得即时化与智能化时，高等教育的育人逻辑又该何去何从？面对颠覆性技术的到来，盲目地拥抱或拒斥技术都无助于解决根本问题，反而容易陷入技术决定论或技术恐惧论的窠臼，必须将生成式人工智能的出现视为倒逼高等教育范式转型的历史拐点。本研究立足我国高等教育知识生产与人才培养的关键命题，旨在回应技术驱动背景下高等教育变革的三个核心问题：第一，生成式人工智能的介入，对我国高等教育的知识范式与育人价值构成了哪些本体挑战？第二，在技术嵌入高等教育的过程中，潜藏着哪些由资本逻辑、文化偏向和数据权力带来的深层风险？第三，面对上述挑战与风险，如何实现从“知识授受”向“原始创新”的根本性跃迁，重塑面向未来、培养拔尖创新人才的中国式高等教育体系？

一、知识范式的根本动摇：生成式人工智能带来的本体挑战

高等教育主要承载着高深知识的保存、传授与发展职能，这一集知识生产与再生产于一体的范式不仅构建了大学作为高深知识殿堂的合法性根基，更维系着社会阶层向上流动的有序通道。然而，生成式人工智能以其颠覆性的知识统整和信息聚合能力，对这一范式的合法性根基发起挑战，这场冲击并非单一维度的技术震荡，而是从培养目标、知识生产过程乃至大学存在价值等本体论层面，对承载着人类智慧传承使命的高等教育体系发起的拷问。实证研究表明，生成式人工智能并非无差别地冲击所有的就业岗位，而是精准地“吞噬”大量初级认知型岗位。相反，需要深度思维与复杂决策的高级岗位的雇佣率不降反增。换言之，生成式人工智能不是取代了人的工作，而是取代了长期以来批量培养的重复性脑力劳动。面对这一变局，当知识的获取不再稀缺，高等教育的育人价值该如何重构？当“认知外包”日益普及，大学学习本身的必要性与价值又该如何确立？更为关键的是，随着知识垄断地位的消解，高等教育“知识本位”范式的内在逻辑又将如何重塑？在此背景下，高等教育的内涵正发生本质嬗变：它不再局限于单向的知识传授与技能训练，而跃迁为一种人机协同下的智慧共生，其核心在于超越对既有知识的

累积与复现，致力于培育能够驾驭算法工具、坚守人类价值的创新人才。这些问题构成了理解并应对高等教育变革的逻辑起点。

（一）知识信仰的解构：高等教育育人价值的合法性危机

人类社会的每一次技术革命，既是生产方式的深刻变革，也是教育逻辑的再造契机。从蒸汽机奏响“大机器生产”的乐章，到电气化拉开新时代的序幕，前两次工业革命的核心逻辑是对“手”（体力劳动）的替代。面对机器换人带来的劳动结构调整，高等教育通过转向对更高级、更复杂的“脑”的专门化训练，确立了其在工业社会不可撼动的地位。与之相对，发端于 21 世纪的信息智能革命，则以知识与算法为中心，开启了对“脑”（脑力劳动）的大规模替代进程。当以 ChatGPT 为代表的生成式人工智能能够胜任文案撰写、编程乃至艺术创作等一系列知识生产工作时，这一图景恰如杰里米·里夫金（Jeremy Rifkin）在其著作《工作的终结》中所预见的：“生产全球所需要的商品和劳务的劳动力却减少了。当前的改革和自动化浪潮只是这一场技术革命的开始，……同时也将使越来越多的劳动力在全球经济中成为多余和无关的人。”这意味着，高等教育赖以存在的“知识－学历－价值”链条正面临前所未有的挑战。

这种替代重心的转移，致使“技术性失业”的风险由体力劳动阶层向知识劳动阶层蔓延。早在 1930 年，经济学家约翰·梅纳德·凯恩斯（John Maynard Keynes）便预言，一种新的疾病在折磨我们，某些读者也许还没有听说过它的名称，不过在今后几年内将听得不想再听，这种病叫作“技术进步导致的失业”。回溯历史，无论是 19 世纪卢德主义者砸向纺织机的铁锤，还是 20 世纪自动化生产线带来的结构性失业，其冲击主要由从事体力劳动的“蓝领阶级”承担。而人工智能时代的“机器换人”正侵蚀知识劳动核心地带。具体而言，企业倾向于大幅缩减对初级知识员工的招聘，并且在初级人才内部，学历背景呈现出 U 型分化，即中层院校毕业生受损最重。当一纸文凭及其所代表的知识劳动不再稳定地通向体面生活，高等教育长期构建的知识改变命运的信仰正被技术悄然消解。

随着技术不断加速知识的生成与传播，高等教育的时间性过程却被不断压缩，学习从“生成意义的过程”异化为“快速获取结果的手段”。加速逻辑在提升知

识流动效率的同时，也消解了高等教育应有的沉潜性，催生出人工智能时代的读书无用论。与以往因回报率低或学用脱节所产生的读书无用论不同，当下的读书无用论是一种更深层次的认知危机和价值危机：当机器在“学习”与“应用”上日益超越人类之时，大学教育存在的意义不言自明，“为什么还要学习”成为时代性的追问。马丁·福特（Martin Ford）敏锐地洞察到，越来越显著的趋势是，许多人即便接受了高等教育，遵循了所有“正确”的路径，却依旧无法在未来的经济形势中找到立足点。种种挑战表明了传统高等教育作为个体安身立命之本和社会阶层跃升之梯的价值正面临重构的压力。

（二）认知过程的悬置：高等教育学习主体的空心化危机

如果说生成式人工智能对高等教育目标的挑战主要体现为知识产出的“可替代性”，那么对教育过程的侵蚀则更为深层且隐蔽，高等教育场域的知识探究本应是一个充满思辨、反思的深层智力活动，但生成式人工智能的介入，使其卷入“加速主义”的逻辑，技术以提升效率为名，将复杂的知识建构压缩为“提问－获取”的瞬时反应。这种“加速”的本质是高等教育的“去过程化”与“去具身化”，它剥夺了学生在知识探究中本应经历的思维磨砺，以及与知识进行深度博弈的认知体验。这使原本鲜活的师生交往被机械地异化为单向度的信息流交互，高等教育过程所独有的生成性、动态性之美，正被算法合乎效率地规训为一种预设的、稳定的标准。

主体性的隐退始于认知外包，其表面的工具理性掩盖了对学术主体性的压缩。英国高等教育政策研究所 2025 年的调查报告显示，生成式人工智能在英国本科生中的学习应用已相当普遍，88% 的学生承认在课业评估中使用 AI 进行辅助（较上年增长 35%）。技术嵌入教育的初衷，本应是降低学习门槛、拓展学习可能性，但在绩点导向与效率至上的双重驱动下，这种辅助极易异化为对思维过程的全面让渡。当学生发现仅凭简化的提示词即可获取结构完整、语言流畅的文本时，那些作为高等教育核心训练的独立思考、文献爬梳与逻辑推演过程便被轻易悬置。此时，技术便从知识探究的“利学”工具异化为“代学”的拐杖，进而诱发认知惰性。当求知的内在渴望因答案的即时获取而日渐萎缩，当学术的尊严因难以甄

别的技术性平庸而受到消解，高等教育面临的将是思维能力的退化与学术精神的空心化。

思维层面的退化，最终将导向哲学层面上的“人的主体性”危机。思想家兰登·温纳（Langdon Winner）曾提醒人类警惕对技术的“反向适应”，如果我们总是透过数据去认识世界，或许会逐步陷入丧失主体性的危机之中，即把由数据所中介化的事实当作客观现实，失去辨识本质领域的冲动与批判质疑的能力，进而迷恋于直观认知，变得懒于思考，从而陷入探及表面真实的虚假满足之中。在传统的高等教育中，学生是学习的主体，知识是客体，技术是中介，然而在“替代学习”的新型人机关系下，这种关系发生了倒置：生成式人工智能凭借其算法和数据成为定义知识、组织答案的“主体”，而学生则退化为提出问题、接收答案的“客体”。在这种主客体关系的倒置中，学习者逐渐失去了在知识建构中的主动地位，其主体意识被压缩为算法系统的输入端。正如尼尔·波兹曼（Neil Postman）所说的，在技术垄断的时代，我们“被解除了思考的责任”，因为“任何技术都能够代替我们思考问题”。当独立思考被算法化思考所取代，当探索的过程被结果的即时化所替代，高等教育便从“立德树人”的灵魂塑造过程，退化为生产“合格使用者”的机制。

（三）教育权威的消解：高等教育范式的秩序重构危机

当高等教育目的与过程被重新定义后，生成式人工智能进一步撼动了“教什么”与“谁来教”的知识根基与权威体系，这预示着高等教育范式面临根本性动摇。首先表现为经典学习理论在高等教育场域的解释力危机。过往，无论是行为主义、认知主义还是建构主义，传统学习理论核心都在于将学习这一人类活动视为可分析、可预测的复杂认知过程。不容忽视的是，当生成式人工智能凭借海量数据训练，能够精准复现“当经验引起个体知识或行动上相对持久的变化时，学习（learning）就发生了”这一经典定义时，一个尖锐的问题便摆在我们面前：沿用了近百年的学习理论，究竟是在描述一种属于“人”的生命实践，还是在定义一种可以被机器模拟甚至超越的信息处理流程？一旦高等教育的学习过程被简化为一种可计算的认知流程，那么机器在逻辑上必然会成为更高效的学习者。

价值根基的动摇，直接体现在教师知识权威的祛魅与角色伦理的重构上。长期以来，教师被视为知识拥有者与传递者，其权威建立在信息稀缺与知识势差的基础之上，然而，当生成式人工智能实现知识的实时生成与迭代优化时，教师不再是唯一的知识中心。面对技术原住民，教师不得不从“知识解释者”转向“学习引导者”和“意义建构的共同体成员”。但这种转型过程充满张力：一方面，部分学生依赖生成式人工智能逃避学习责任，甚至以算法生成的内容挑战教师判断；另一方面，教师权威陷入既要维系教育的认知深度，又要守护教育人文精神的内在张力与角色冲突之中。教师权威与教学范式的双重动摇，标志着传统“知识本位”教育体系正陷入范式困境。这让我们不得不重审托马斯·库恩（Thomas Kuhn）的“范式革命”理论。长期以来，“教师中心、书本中心、课堂中心”以及“讲授－接收”的单向范式，构成了大学最为稳固的教学图景。但当知识的获取不再受制于大学围墙，当教师不再享有知识传播的垄断权时，这套范式的合法性便受到了根本性质疑。钱颖一教授曾指出，中国教育的系统性偏差在于把教育等同于知识并局限于知识。生成式人工智能的出现，清晰地照出“知识本位”范式早已存在的内在矛盾。据此，高等教育的根本危机并非源自外部技术，而是源自教育自身对“知识传递”功能的过度依赖与对“立德树人”本质的偏离。在这一意义上，生成式人工智能的挑战恰恰构成了“不破不立”的契机，它倒逼高等教育必须重新审视如何在技术理性与人文理性之间重建平衡，如何在算法效率与育人价值之间划清边界。

二、外部逻辑的深度介入：生成式人工智能介入的深层风险

本体挑战的出现并非单纯的技术迭代现象，而是更为深刻的社会结构性力量在高等教育领域的集中投射。正如 OpenAI 公司的首席执行官山姆·奥特曼（Sam Altman）所警示的：人类的未来要由人类自己决定。这一论断在高等教育语境下的隐喻为：大学作为人类理性与良知的精神堡垒，如果轻率地将关乎知识生产与人才培养的核心权力让渡给由商业逻辑驱动的智能系统，那么高等教育将面临从办学自主权、知识定义权到国家意识形态安全的风险。因此，在将生成式人工智能作为辅助工具引入高等教育场域之前，必须对其背后的资本逻辑与潜在的后果进行反思与评估。

（一）资本逻辑的渗透：高等教育公平的异化

理论上，生成式人工智能具备打破知识传播时空壁垒的潜力，为通过高等教育实现阶层跃升提供了新的契机，这也正是技术乐观主义者所描绘的“推倒大学围墙”后的普惠图景。然而，技术的实际演进并非由纯粹的公益理想所主导，科技公司意在通过垄断算力与数据优势来攫取超额利润，这使得技术的应用重心天然地偏向于市场效益而非教育价值。马丁·海德格尔（Martin Heidegger）强调，现代技术从诞生之初，其发展逻辑便与资本的逐利本性紧密捆绑。当以 OpenAI 为代表的科技巨头迅速完成对知识生产工具的资本化闭环时，大学所坚守的公共性与资本固有的逐利性之间的张力，便在高等教育场域中日益凸显。

这种张力，使得技术赋能高等教育的“普惠性”承诺在现实中不可避免地走向异化。资本常通过“免费增值”的商业模式介入教育，制造出人人皆可触及的普惠表象。然而，无论是更深度的逻辑推理模型、更前沿的学术数据库，还是定制化的科研辅助服务，都被悉数置于高昂的费用之后。吴康宁指出，人工智能的条件差异会进一步加剧因经济基础差异所导致的教育机会不平等。在这一逻辑下，所谓的“技术公平”荡然无存，高等教育也在此过程中，从维系社会流动的公共阶梯，变成一种价高者得的营利商品。

随之而来的是，一种比传统的物理资源不均更隐蔽、更难跨越的智能鸿沟正在大学校园内生成。来自优势阶层的学生有能力购买顶级的 AI 服务，将其转化为学术论文润色、留学申请优化乃至复杂编程辅助的优势；而来自农村或贫困家庭的学生，往往受限于经济能力，只能使用功能受限的基础版本，甚至因缺乏 AI 素养而沦为算法的被动消费者。这道由资本划定的“算法鸿沟”与“算力鸿沟”，使得技术非但未能成为弥合教育差距的桥梁，反而可能沦为固化甚至加速精英阶层优势积累的工具。由此，高等教育的马太效应在智能时代不仅未能消解，反而披上了更隐秘、更强大的技术外衣。

（二）技术理性的扩张：高等教育价值的异化

高等教育的根本任务在于“立德树人”，其终极旨归是培养德智体美劳全面

发展的社会主义建设者和接班人。然而，当高等教育在“知识本位”范式的动摇中陷入价值迷茫时，技术理性便在大学价值真空中长驱直入，致使大学的价值取向正面临被窄化的风险。“技术中性论”认为技术本身没有价值观，其价值取向取决于使用者，但这一论断在生成式人工智能时代尤为需要警惕，因为技术并非价值的真空地带，生成式人工智能的“智能”源于对海量互联网文本数据的学习，而这些数据本身就充满了特定文化背景下的偏见、立场与价值判断。当学生将其作为知识探究的主要依托时，便极易被算法的隐性偏见所裹挟，导致学术判断力的丧失与价值立场的虚无。

其危险性在于，技术的效率至上正将高等教育的价值追求引向一种可计算、可量化的单维方向。在这种范式下，技术的元价值观——效率、竞争和产出——被奉为大学治理的圭臬。高等教育的成功标准不再锚定于培养具有家国情怀与创新精神的栋梁，而被异化为对绩效、排名、升学率等可测量指标的极致追逐。诚如卢梭曾论述的，“随着科学与艺术的光芒在我们的地平线上升起，德行也就消逝了；并且这一现象是在各个时代和各个地方都可以观察到的”。这一历史的忧思在当今表现得更为具体：当知识生产被算法辅助主导，当教育成效皆被数据量化时，那些无法被算法编码的人格养成、批判性思维与社会责任感便被系统性地边缘化。这恰恰印证了钱穆先生的忧思：“教人作人，亦分内外两面。知识技能在外，心情德性在内。做人条件，内部的心情德性，更重要过外面的知识技能。”生成式人工智能可以高效地生产“外在”的学术文本与技能，却难以涵养“内在”的学术良知与德行。

最终，这种价值的单维化将导致高等教育“育人”使命的全面消解。当“有效性”与“效率”成为大学的唯一标尺时，教育目的将不可避免地被窄化为智育的单一培养。大学不再追问“为何学”的终极价值，而只关心“学什么能带来最大收益”的工具价值。在这种技术化的生存方式中，人逐渐忘却了自身的生存意义，思想家埃·弗洛姆（Erich Fromm）强调，“人创造了种种新的、更好的方法以征服自然，但他却陷入这些方法的网罗中，并最终失去了赋予这些方法以意义的人自己。人征服了自然，却成了自己所创造的机器的奴隶”。当大学生沉浸于与机器的高

效交互，却逐渐丧失与导师、同窗等进行真实思想碰撞与情感交流的意愿时，大学作为精神共同体的属性荡然无存。高等教育从“育人”（培养具有主体性的完整的人）滑向“制器”（生产适应系统的精巧工具），这构成了人工智能时代高等教育最深层的价值危机。

（三）数据权力的垄断：高等教育认知的异化

随着人工智能在教育领域的深度嵌入，数据的价值超越了技术要素本身，跃升为关乎国家竞争力的核心战略资源，究其本质是认知主权的竞争。早在信息时代初露端倪时，思想家尼尔·波兹曼（Neil Postman）就曾发出警示，“电脑对他们有多大的好处呢？他们的隐私更容易被强大的机构盗取。他们更容易被人追踪搜寻、被人控制，更容易受到更多的审查，他们对有关自己的决策日益感到困惑不解；他们常常沦为被人操纵的数字客体。他们在泛滥成灾的垃圾邮件里苦苦挣扎。他们容易成为广告商和政治组织猎取的对象”。这一断言在智能时代具有了更为紧迫的意义：当一个国家未来精英的知识获取路径、思维范式乃至价值观的生成过程，深度依赖于少数不受有效监管的，甚至由他国资本掌控的科技巨头所提供的算法黑箱时，高等教育面临的已不仅仅是技术层面的数据安全问题，更是国家教育主权与意识形态安全虚化的风险。这种虚化并非一蹴而就的显性对抗，而是通过两种隐蔽的认知规训，潜移默化地侵蚀着高等教育的自主性。

首先是算法精心编织的“信息茧房”。在平台算法的持续追踪与精准的用户画像机制支持下，生成式人工智能所产出的内容被定向筛选并持续推送，以迎合用户既有的立场与认知舒适区。这种机制将正处于思维活跃期的青年学子桎梏于一个看似广博、实则封闭的茧房之中。在高等教育最为珍视的批判性思维培养中，学生接触多元异质观点的机会被个性化推荐算法过滤，思想碰撞的火花被数据平抑，批判性思维因缺乏对立面的砥砺而日渐钝化。长此以往，未来的学术精英将面临认知视野的窄化与思维模式的单向度化，最终丧失对复杂世界进行全面、辩证和深度认识的能力。其次是模型自身固有的“认知幻觉”对学术求真精神的消解。生成式人工智能因其技术原理限制，并不真正理解语义和真实世界，而是基于概率生成文本。这导致它会以极其自信、权威的语气，生成看似逻辑严密却完全基

于虚构的信息、事实乃至引文，即所谓的“机器幻觉”。研究揭示，即使是性能优越的 ChatGPT-4，在处理高等数学（如微积分）等严谨的逻辑问题时也表现不佳。这种“权威的谎言”比单纯的错误知识更具危害性，因为它极大地增加了辨别成本，侵蚀了教育信任的基础。更深层的隐患在于，由算法主导的“信息茧房”与“认知幻觉”，共同构成了对国家教育主权虚化的深层威胁。当高等教育中知识的流通、定义的逻辑由外部商业平台乃至他国意识形态所隐性控制时，国家意识形态的屏障将面临风险。面对这一风险，中国高等教育的核心命题不再仅仅是如何利用技术提升效率，而是如何在技术洪流中坚守“为谁培养人”的政治方向，守护国家教育的价值主权与文化安全。

三、走向创新本位：高等教育育人体系的根本性重构

面对生成式人工智能技术带来的挑战，无论是选择性地限制、封堵，还是主动地适应、变革，都考验着我国高等教育体系的战略智慧与历史担当。若说传统高等教育主要承载着对人类既有存量知识进行编码与传递的职能，旨在培养学识渊博的“百科全书式”人才，那么生成式人工智能的崛起则以其超越性的知识整合与生成能力，揭示了高等教育的局限性。当机器在记忆、检索和组织知识的效率上远超人类时，“大学何为”以及“何以为学”重新成为必须直面的命题。由此，高等教育的根本出路已无法依靠单纯的技术叠加或局部的课程修补，而必须触及学术生产的本源问题，实现从知识授受向知识创造的本质性价值转向与逻辑重构。

（一）重塑学习目的：从“工具理性”走向“价值理性”

长期以来，高等教育场域在实践中不仅面临“知识本位”的固化，更被外在的功利主义所裹挟，学生的学习目标往往被窄化为通过刷绩点获取保研资格、考取光鲜文凭、谋求高薪职位等短期工具性目标，这种以工具理性为主导的教育观，使得原本神圣的求知探索过程异化为一场追求外部回报的符号交换，而非生命的内在成长。管理学大师彼得·德鲁克（Peter Drucker）曾用“三个石匠”的故事来比喻不同层次的工作动机：第一个石匠为养家糊口而工作，第二个为成为最好的工匠而工作，第三个则为建造一座宏伟的教堂而工作。同理，如果大学学习仅仅

是为了“绩点至上”的短期功利目的，那么当生成式人工智能能够比大多数学生更高效地撰写报告、优化代码，甚至通过标准化考核时，这种单一的功利主义路径就会变得愈发狭窄和脆弱。

因此，应对生成式人工智能挑战的首要任务，便是引导整个高等教育体系实现育人目标的根本性重塑，即完成一场大学精神的价值理性回归，这要求高等教育必须破除“唯分数、唯升学”的顽疾，引导学生从“基于无知而求知”转向“为了超越而创造”，在新的时代条件下重新思考“为谁培养人”这一根本政治命题。正如克里希那穆提所言：“技术永远无法产生创造性的了解。……真正的教育，乃是帮助个人，使其成熟、自由，绽放于爱与善良之中。”人工智能的强大恰恰在于处理“外在”的知识与技能，而高等教育的不可取代性则在于培育无法被算法编码的“内在”的精神与德性。

重塑学习目标，关键在于将高等教育的聚焦点从“既有范式的复现与应用”转向“未知领域的探索与突破”，从内心深处激发并守护学生与生俱来且无法被算法编码的两大根本能力：一是超越知识本身的好奇心与想象力。阿尔伯特·爱因斯坦曾说过，想象力比知识更重要。这是因为知识是有限的，而想象力却包含整个世界。以标准化考核为核心的传统高等教育模式，在追求标准的过程中往往扼杀了学生的知识探究与创新能力，在生成式人工智能时代，无论是本科教学还是研究生培养，大学都应致力于为学生创造能够大胆试错、自由探索的环境。一如陶行知所主张，教育不能创造什么，但它能启发并解放儿童的创造力，让他们去从事创造性的工作。二是深刻的批判性反思与提出真问题的能力。人工智能可以基于海量数据回答“从1到10”的问题，但它无法像人类学者一样，基于对世界和自身的深刻体悟，提出“从0到1”的创见。在一个人人都能轻易获得“答案”的时代，提出好问题比拥有答案更稀缺、更有价值。因此，高等教育必须鼓励学生敢于质疑权威、善于提出真问题，唯有如此，才能培养出真正能够驾驭技术、引领发展而不是被技术洪流所裹挟和替代的拔尖创新人才。

（二）变革学习范式：从“存量知识的再生产”走向“增量知识的创生”

工业社会以来，高等教育实践被量产模式主导，其核心是建立一种“面向已

知世界”的学习范式。在这种认知闭环中，学生的认知逻辑遵循着“知识－问题－方法－知识”的闭环路径，这种路径在工业时代无疑是高效的，因为它旨在批量培养能够精准执行重复性任务的标准化人才。而在生成式人工智能时代，传统范式的基石已被动摇。人工智能的强大之处在于，它能够比人类更出色地“计算”并完成“应用所知、复现已知”的任务。如果高等教育继续停留在对存量知识的机械记忆上，无异于将人类置于与机器进行同质化算力竞争的必败窘境。

因此，高等教育的当务之急，是引领一场从“面向已知”到“面向未知”的认识论革命。这种创新性学习以培养人的高阶思维为宗旨，其核心是从静态的“知识积累”走向动态的“新知创生”。其逻辑轨迹不再是封闭的循环，而是开放的、螺旋式上升的，遵循着“真实问题——演进逻辑——学术前沿——多维方法——原始创新”的路径。这是一种拥抱不确定性、鼓励认知冒险的学习，其所追求的并非效率最优的标准解，而是最具主体创造性的“非标准答案”。

这种学习范式的变革，本质上是激活并释放人类独有的“算计”智慧，即面向未来的、包含价值判断与复杂决策的高级智能活动。它由一个以学生为主体的“四轴驱动”模式所牵引：一是以“好奇心”驱动“真问题”的发现。在生成式人工智能时代，高等教育的首要任务是将学习的重心从“求解”转变为“提问”，引导学生从原本碎片化的网络知识中跳脱出来，在真实的学术与社会场景中发现那些 AI 无法定义的“真问题”。二是以“批判性”驱动“认知自主”的实现。生成式人工智能提供的内容往往缺乏对本质的深刻洞见，这反而为批判性思维的训练提供了前所未有的“靶场”，创新性学习要求学生不唯书、不唯 AI、只唯实，在与 AI 的辩证博弈中实现对既有理论的反思和超越。三是以“想象力”驱动“大概念”的整合。AI 擅长处理单一维度的任务，而人类擅长跨界联通。高等教育应鼓励学生打破学科壁垒，将孤立的知识点串联成线、编织成网，在看似无关的领域间构想新的可能，培养复合型的创新思维。四是以“责任心”驱动“价值理性”的回归。知识的创造不应是象牙塔内的智力游戏，而应服务于人类福祉。AI 只能处理实然的事实，无法处理应然的价值，这正是高等教育“立德树人”的独特价值，通过“经艰难事”的磨砺，培养学生深厚的家国情怀与社会责任感，确保技术的力量始终掌握在具有道德自觉的人的手中。

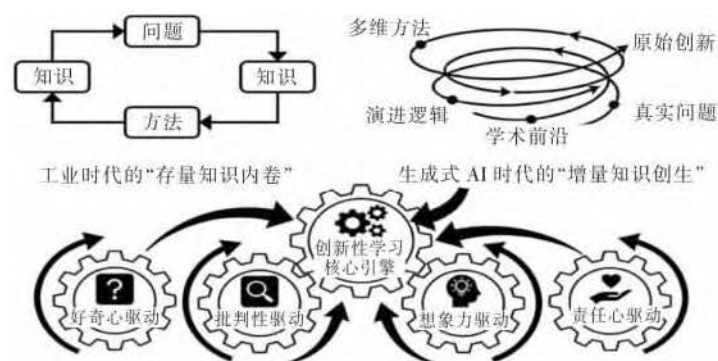


图1 “面向未知世界”的创新性学习体系

（三）再造评价体系：从“量化指标的崇拜”走向“育人本位的回归”

评价体系不仅是教育质量的标尺，更是学生身心成长与价值塑造的指挥棒。中共中央、国务院印发的《深化新时代教育评价改革总体方案》明确提出，要“坚决克服唯分数、唯升学、唯文凭、唯论文、唯帽子的顽瘴痼疾”。这一国家意志在人工智能时代显得尤为紧迫。长期以来，高等教育评价中盛行“绩点为王”与“量化崇拜”，将学生丰富立体的生命成长异化为冷冰冰的数字堆叠。这种评价逻辑导致学生陷入功利主义泥潭，不仅窄化了育人的内涵，更造成了学生精神世界的空心化。在生成式人工智能时代，这种评价体系的危机被彻底引爆：当算法能够以极低成本批量生产标准答案与规范文本时，若高等教育仍将评价锚定于知识复现与文本堆砌，那么学生注定沦为劣势一方。因此，为了捍卫高等教育培养拔尖创新人才的使命，评价体系必须完成从“工具指标”向“育人本位”的转型。这场转型的核心，在于跳出单纯的产出导向，建立一套能够确证人类主体价值、彰显道德品格与综合素养的评价新范式。

第一，重构评价内容，引导学生从“量化产出”到“素养生成”的价值转向。面对AI带来的知识通胀，新的评价体系应超越单纯的学业成绩与技能产出，转而关注那些无法被算法编码的人的特质。其一，确立“德行”为首要尺度。在人机共生的社会，比智力更重要的是驾驭技术的伦理自觉。评价不应局限于课堂，而应重点考查学生的数字伦理、社会责任感以及在复杂情境中的价值判断能力。其二，确立“多元”的成才观。对于本科生而言，应评价其是否具备独立思考的批判性思维、跨学科解决问题的实践能力以及终身学习的元认知能力。对于研究生，

则侧重评价其原始创新能力与学术伦理。无论是投身社会实践、从事艺术创作，还是解决工程难题，只要体现了人的创造性与主体性，都应被视为有价值的成长，进而引导学生从“与机器比拼知识存量”的困境中走出来，转向追求自身精神丰盈与人格完善。

第二，创新评价方式，实现从“单一结果甄别”到“全过程生命成长”的范式进阶。教育评价的本质不是对学生进行分层，而是支持其发展。变革评价方式，意味着要建立起一套能够容纳差异、激励个性的全过程伴随式评价框架。一方面，推动评价场景的“真实化”。将评价从封闭的教室中解放出来，投射到真实的社会生活与实践场域中。鼓励学生在田野调查、志愿服务、企业实习等真实情境中去发现问题、解决问题，将其在实践中展现的协作精神、抗逆力与同理心作为评价的重要依据。这不仅适用于学术研究，更是通识教育与职业素养评价的题中应有之义。另一方面，推动评价手段的数智化赋能。智慧地利用技术构建伴随式的成长反馈系统，但其目的不是监控，而是看见与支持。通过捕捉学生在项目协作、试错迭代过程中的思维路径与行为数据，生成多维度的成长画像，帮助学生在与AI的交互中清晰界定自身的优势与短板，从而实现从单纯关注冰冷的“分数结果”转向关注有温度的“生命成长”。

作者信息

邬志辉 东北师范大学中国农村教育发展研究院教授、博士生导师

龚毓琳 东北师范大学中国农村教育发展研究院博士研究生

王小飞

联合国教科文组织人工智能教育立场的
阐释与实践启示

来源 | 《教育国际交流》2026 年 03 期



中国教育科学研究院比较教育研究所所长、研究员 王小飞

生成式人工智能的爆发式发展，正以颠覆性力量重塑全球教育生态，同时也对全球人才结构与能力体系提出了前所未有的挑战。世界经济论坛的研究显示，到 2030 年，全球近 60% 的劳动者将需要接受技能重塑与更新，AI 时代的人才培养已成为各国教育变革与国家竞争力构建的核心议题。在此背景下，联合国教科文组织（UNESCO）在 2025 年 9 月发布的报告论文集《人工智能与教育的未来：变革、困境与方向》（以下简称“文集”）中，汇聚全球高端智库、知名教育学者、高新企业高管等各方多元观点，聚焦教育在人工智能时代所面临的问题、两难境况及发展趋势，开展了全面深入的研讨。该文集深入哲学、伦理、教育学与社会学维度，审慎对待关于 AI 教育的“乌托邦”式愿景和“反乌托邦”式恐惧，将其定位为需要被规范、引导与重塑的社会技术系统，从而为 AI 时代的全球人才培养划定价值底线与实践方向。

UNESCO 关于 AI 教育的立场，本质上是对教育本质的坚守与对技术理性的反思。其核心观点在于：人工智能是服务于教育公平与人的全面发展的工具，而非替代教育人文性、社会性与关系性；AI 融入教育必须以人本为基础、以公平为导向、以伦理为前提、以人文为内核，坚决抵制商业资本对教育公共性的侵蚀，反对算法霸权对教育多样性的消解，推动构建公平、包容、人文导向的 AI 教育新生态。本文基于该文集的核心内容，从其价值取向、教学主张、学习观点、评价主张、治理框架等维度，阐释 UNESCO 关于 AI 时代人才培养的基本主张，并结合我国当前 AI 教育实践，提出相应的思考建议。

一、AI 时代人才培养的价值根基：人文主义立场

UNESCO 关于 AI 教育的所有主张，均建立在其人文主义立场之上，是其区别于商业机构、技术企业与单一国家 AI 教育理念的核心特征。其价值取向主要集中于教育的公共性、公平性、人文性以及包容性等维度，秉持非技术中心、非商业主导的 AI 教育主张。

（一）坚守人才培养的公共性本质，反对商业资本主导

UNESCO 明确将教育定义为全球公共产品，人才培养作为教育的核心目标，同样必须坚守公共属性。文集在前言中指出：AI 融入教育不能成为商业资本逐利的工具，不能将公共教育体系重塑为私人资本的盈利场域。当前全球 AI 教育领域，少数商业巨头掌控前沿模型研发与推广，以效率、规模、优化为核心叙事，挤压教育的公共空间，忽视教师、学习者与社区的核心诉求。UNESCO 坚决反对商业导向的 AI 教育变革，主张 AI 技术的研发与应用必须服务于教育公共利益，由公共部门主导、多元主体参与构建全球对话共同体，确保 AI 教育始终服务于全民终身学习、包容公平教育的 2030 教育目标，而非商业企业的市场扩张与资本增值。

（二）以公平包容为核心，消解 AI 带来的教育不平等

公平与包容是 UNESCO 关于 AI 教育价值取向的核心要义。文集第七章“直面教育中的隐性不平等”集中阐述了 AI 技术客观上有可能带来的不平等现象，指出其应用加剧了全球教育的数字鸿沟、语言鸿沟、性别鸿沟与资源鸿沟：发达国

家与发展中国家、富裕群体与贫困群体、主流语言群体与小众语言群体，在 AI 人才培养资源获取、使用与受益等方面存在巨大差距。这种差距不仅是技术接入的差异，更决定了不同群体能否平等参与人才培养，能否获得进入未来劳动力市场的核心能力。为此，UNESCO 秉持其一贯的公平、包容的立场，主张优先保障边缘群体、弱势群体、欠发达地区的 AI 教育权益，推动多语言 AI 系统研发、本土文化适配的 AI 设计，反对 AI 成为强化教育不平等、固化阶层差异的工具，致力于通过 AI 技术缩小而非扩大教育差距。

（三）坚守人才培养的人文性内核，捍卫教育的关系性与社会性

教育的本质是关系性、社会性与伦理性的实践，而非单纯的知识传递与技能培养。AI 技术的算法化、高效化、个性化特征，容易消解教育中慢思考、深反思、人际连接的核心价值，将教育简化为数据处理与认知输出。文集中有多位论者主张，人才培养的核心是人的成长、人际关怀、价值塑造与社会联结，人工智能无法替代教师的情感关怀、师生的情感互动、同伴的协作交流与文化的传承创新。因此，UNESCO 主张 AI 教育必须坚守人文性内核，将人的尊严、团结、正义与关怀置于核心位置，反对技术对教育人文性的侵蚀，避免 AI 将教育异化为标准化、算法化、非人性化的流水线，始终捍卫教育的社会属性与人文温度。

（四）尊重多元文化与知识体系，反对技术霸权与知识单一化

UNESCO 反对 AI 人才培养中的技术霸权与知识单一化，主张尊重全球多元文化、本土知识、小众语言与传统教育智慧，培养具有文化多样性与本土适应性的人才。当前主流 AI 模型多以西方文化、英语语言、现代知识体系为基础，忽视全球多元文化与本土知识的价值，导致教育中的文化同质化与知识霸权问题。如在该文集所收录的巴约·阿科莫拉菲访谈录中，巴约指出：“我们再也不能以我们习惯的、由欧美启蒙运动所倡导的方式来思考自己，……这种做法会遮蔽世界运转和自我生成的方式。”UNESCO 通过该文集强调，AI 教育必须尊重不同国家、不同民族、不同社群的文化特性与知识体系，推动 AI 技术的本土化、语境化适配，支持多语言、多文化、本土知识导向的 AI 系统研发，反对单一文化、单一知识体系通过 AI 技术主导全球人才培养，维护全球人才培养的文化多样性与知识多元性。

二、AI 时代人才培养中的教师：人的不可替代性

在教学实践层面，UNESCO 立足人文主义教育观，提出坚守教师核心地位，强调教育优先、关怀导向等方面主张，反对技术替代教师、算法主导教学的错误倾向。

（一）坚守教师的核心地位，明确 AI 的辅助角色

UNESCO 在其发布的多篇报告中反复强调，人工智能永远无法替代人类教师，AI 只是服务于教学的辅助工具。教育是培养完整的人的实践，需要教师的情感关怀、价值引领、道德示范与个性化引导，这些是 AI 无法具备的人文属性。UNESCO 主张，教师在 AI 教学场景中是学习的设计者、伦理的守护者、关系的构建者、价值的引领者，而 AI 则承担个性化辅导、重复性工作、资源供给、数据反馈等辅助性任务，形成“教师主导、AI 辅助”的教学格局。

与此同时，UNESCO 主张，AI 时代的教师必须具备专业的 AI 素养与伦理引领能力，这是保障 AI 教学良性发展的核心。文集所收录的“在 AI 时代坚守教育目标”一文中指出：需要从关系、目的、认知、心理以及教学等五个维度理解 AI 对教育的影响，坚守教育的育人本质，抵制 AI 对学生自主性、好奇心与情感健康的侵蚀。此外，UNESCO 在《教师 AI 素养框架》中也明确指出：教师需要掌握 AI 伦理判断能力，引导学生正确使用 AI 技术，识别算法偏见、虚假信息与伦理风险，成为 AI 教育的伦理守护者。反对将教师简化为 AI 工具的操作者，主张培养教师成为 AI 教学的设计者、监督者与引领者，通过参与式设计、伦理审计、数据治理等方式，主导 AI 技术在教学中的应用。

（二）推行教育原则优先于技术的人才培养逻辑

UNESCO 反对“技术先行”的教学变革模式，主张教学实践必须坚守教学优先原则，技术服务于教学目标而非主导教学方向。文集指出，当前全球 AI 教育实践普遍存在“技术至上”误区，学校盲目引入 AI 工具，却忽视教学目标、学生发展与教育规律，导致技术与教学脱节。UNESCO 主张，AI 融入教学必须先明确教育目标与教学理念，再选择适配的 AI 技术，应始终以教学规律、学生身心发展规

律为核心，避免技术绑架教学、算法替代教学的现象。无论是 AI 辅助教学、个性化学习还是智能评价，都必须以育人为基础，技术只是实现教学目标的手段，而非教学的核心。

（三）推行关怀导向的教学设计，嵌入伦理与人文关怀

基于人文关怀和伦理考量，UNESCO 主张 AI 时代的教学必须嵌入关怀导向的设计，将人文关怀、情感支持、公平包容融入 AI 教学全流程。文集中提出“关怀设计”七大转变：从技术主导转向师生参与式设计、从效率导向转向关怀导向、从数据监控转向信任构建、从算法偏见转向公平导向、从标准化转向个性化关怀、从商业主导转向公共利益、从技术替代转向人机协同。这种教学设计强调，AI 教学必须关注学生的情感需求、心理发展与社会归属，教师通过 AI 工具更好地实施个性化关怀，要始终以学生的全面发展为核心，避免算法对学生的标签化、歧视化处理。

三、AI 时代人才培养的核心重构：面向人机共生的能力体系

AI 技术的发展正在重构人才的能力需求，UNESCO 立足学生主体性与全面发展理念，提出 AI 时代人才培养必须重构核心能力体系，重点培养学生的自主学习能力、AI 核心素养、社会化协作能力，反对认知卸载、被动接受、算法控制与不平等学习，构建人机共生的人才培养模式。

（一）坚守学生的学习主体性，抵制认知卸载与思维退化

UNESCO 教育助理总干事斯蒂芬妮亚·贾尼尼在该文集前言中指出：教育日益成为商业化人工智能模型的前沿战场，这些模型承诺带来速度、效率和规模，同时也威胁到深度学习，有可能造成认知卸载、批判性思维下降。UNESCO 明确反对 AI 技术导致的学生认知卸载与批判性思维退化，强调学习的本质是学生主动建构知识、发展思维、提升能力的过程，AI 不能替代学生的主动思考与深度认知。文集强调，生成式人工智能能够快速生成答案、完成作业，容易导致学生产生认知惰性，过度依赖 AI 工具，丧失独立思考、问题解决与创新能力。同时，AI 系统基于历史数据训练，具有“向后看”的特征，容易固化学生的思维模式，消解

创新与批判精神。为此，UNESCO 主张，必须坚守学生的学习主体性，引导学生将 AI 作为学习工具而非替代者，培养学生的主动探究、深度思考与批判认知能力，坚决抵制 AI 导致的思维退化与认知懒惰。

（二）培养学生的 AI 时代的核心素养，构建人机共生的基础能力

AI 素养、批判性思维、创新思维与伦理认知应该成为 AI 时代学生的核心素养。安德烈亚斯·霍恩在文集中系统阐述了 AI 时代学生所需具备的概念素养、批判素养与创造素养三大核心能力。其中，概念素养要求学生理解 AI 的基本原理、算法逻辑与技术局限；批判素养要求学生识别 AI 的偏见、虚假信息、伦理风险与权力关系；创造素养要求学生合理、合规、创造性地使用 AI 技术。AI 素养不是可选的附加素养，而是与读写算同等重要的基础素养，必须纳入全学段教育体系，尤其保障边缘群体、弱势群体的 AI 素养教育，避免数字文盲加剧教育不平等。学生必须掌握 AI 伦理规范，树立数据隐私意识、责任意识、公平意识，成为负责任的 AI 使用者。有鉴于此，UNESCO 主张，AI 时代的人才培养需要超越传统的“智能范式”，转向“智慧范式”。北大博古睿研究中心宋冰教授在该文集中指出：“当前的人才培养过度聚焦于技能与知识的传递，而 AI 时代更需要培养学生的智慧，包括对自我的反思、对世界的认知、伦理判断与价值选择，这是应对 AI 时代不确定性的核心能力。”这样的智慧教育，要求人才培养不仅要让学生掌握使用 AI 的技能，更要让学生学会反思技术的影响，坚守人的主体性，实现人与 AI 的协同共生。

（三）反对过度个性化，倡导社会化学习与集体智慧

UNESCO 对 AI 驱动的超个性化学习持批判态度，主张学习是社会性实践，反对算法个性化导致的学生孤立化与“回音室”效应。保罗·普林斯洛在该文集所收录的“幼稚化、回音室效应、过滤泡沫，还是新启蒙的曙光”一文中强调：AI 驱动的个性化学习虽然能够适配学生的学习节奏与需求，但过度个性化会隔离学生与同伴的交流，缩小学生的认知视野，强化固有认知，形成“学习回音室”，消解学习的社会性与协作性。有鉴于此，AI 时代的学习必须平衡个性化与社会化，推动再社会化学习，将 AI 作为集体智慧、协作学习的支持工具，促进学生之间的

交流互动、合作探究，培养集体智慧与社会交往能力，避免技术将学习异化为孤立的个体行为。

（四）保障所有学生的公平学习权，消解人才培养鸿沟

公平学习是 UNESCO 关于学生学习的核心价值取向，应保障所有学生平等享有 AI 学习的机会、资源与能力，消解 AI 带来的人才培养不平等。该文集注意到当前 AI 学习存在严重的不平等现象，如发达国家学生能够使用前沿 AI 工具，发展中国家学生缺乏基本的网络与设备 □ 主流语言学生能够无障碍使用 AI，小众语言学生面临语言壁垒 □ 富裕家庭学生能够享受个性化 AI 学习服务，贫困家庭学生被排斥在外。UNESCO 主张优先保障欠发达地区、边缘群体、小众语言群体、残障群体的 AI 学习权益，推动离线 AI 工具、多语言 AI 系统、低成本 AI 设备的研发与应用，将 AI 作为促进学习公平的工具，而非扩大差距的壁垒。

（五）尊重学生的认知自主与数据主权，保护学习隐私

UNESCO 高度重视学生的学习隐私、数据主权与认知自主，反对 AI 学习中的数据滥用、监控过度与认知控制。文集中指出，AI 学习系统收集大量学生的学习数据、行为数据、心理数据，存在数据泄露、隐私侵犯、算法监控、认知操控等风险，侵犯学生的人格尊严与数据主权。同时，AI 系统的算法推荐可能控制学生的学习内容与认知方向，消解学生的认知自主与选择自由。为此，UNESCO 主张，必须严格保护学生的学习隐私与数据主权，明确数据收集、使用、存储的伦理边界，推行学生数据的教师主导治理、参与式监管，反对商业机构滥用学生数据牟利，保障学生的认知自主与学习自由。

四、促进 AI 时代教育评价创新：去标准化、重过程、守真实、促公平

UNESCO 直面生成式人工智能对传统评价体系的冲击，提出去标准化、重过程、守真实、促公平等主张，重构 AI 时代的教育评价逻辑，反对 AI 虚假评价与评价不平等。

（一）批判传统标准化评价，倡导人才评价体系革新

生成式人工智能的出现暴露了传统标准化评价体系的致命缺陷，传统的纸笔测试、论文写作、作业评价等方式，无法区分学生是自主完成学习任务还是借助 AI 生成内容，导致评价的真实性、有效性大幅下降。传统标准化评价注重知识记忆、标准化答案，难以评价学生的批判思维、创新能力、协作能力与问题解决能力。迈克·珀金斯和贾斯珀·罗在该文集所收录的“我们所熟知的评估时代的终结”一文中指出：“生成式人工智能的出现可能导致传统评估方式的崩溃，但也为我们提供了机会，让我们仔细思考未来的评估方式会是什么样子。”由此，UNESCO 强调，应推动评价体系的系统性革新，从知识本位评价转向素养本位评价，从结果本位评价转向过程本位评价，从标准化评价转向多元化评价，引导人才培养聚焦 AI 时代核心能力的提升。

（二）推行过程性、形成性评价，弱化终结性评价

AI 时代的教育评价应强化过程性、形成性评价，弱化终结性、标准化评价，发挥 AI 技术在过程性评价中的优势。UNESCO 在文集中提出，AI 技术能够实时收集学生的学习过程数据，提供及时、精准的形成性反馈，辅助教师开展过程性评价，帮助学生了解学习进展、调整学习策略。这种评价方式符合学习的发展性规律，能够更好地促进学生学习，替代传统的终结性应试评价。因此，有必要利用 AI 技术构建连续、动态、人性化的形成性评价体系，将评价融入学习全过程，评价目的从甄别选拔转向发展促进，发挥评价的育人功能而非筛选功能。

（三）抵制评价不平等，构建公平包容的 AI 人才评价体系

UNESCO 将公平性贯穿 AI 教育评价始终，反对 AI 评价中的不平等现象，主张构建公平、包容、无偏见的评价体系。UNESCO 在该文集中指出，AI 评价系统基于主流文化、主流语言、优势群体数据训练，存在文化偏见、语言偏见、性别偏见与阶层偏见，对边缘群体、小众语言群体、弱势群体的评价存在歧视，加剧评价不平等。同时，发达地区与欠发达地区在 AI 评价资源上的差距，导致评价机会的不平等。为此，必须对 AI 评价系统进行偏见审计、公平性审查，推动多语言、

多文化、本土适配的评价系统研发，保障所有学生在评价中的公平权益，避免 AI 评价成为固化教育不平等的工具。

五、AI 时代人才培养的制度保障：多元协同的治理体系

在教育治理层面，UNESCO 提出全球协同、多元参与、伦理先行、法治保障、本土适配的核心主张，构建政府、国际组织、学校、教师、学生、社区共同参与的 AI 教育治理体系，反对技术霸权、商业垄断与单一治理模式。

（一）坚持多元主体参与，构建民主共治格局

UNESCO 倡导多元主体参与的 AI 教育治理格局，反对政府单一管控或商业资本主导治理，主张政府、国际组织、学校、教师、学生、家长、社区、技术企业共同参与治理。文集中强调，AI 教育治理不能由技术企业或少数精英主导，必须保障教师、学习者、边缘群体的话语权与参与权，尤其是一线教师、弱势群体、本土社群的声音必须被纳入治理决策。主张在学校层面，推行教师与学生参与的 AI 治理机制；在国家层面，构建政府统筹、多方参与的治理体系；在全球层面，强化国际组织的协调作用，形成民主、包容、多元的共治格局，确保 AI 教育治理符合公共利益与教育本质。

（二）强化政府公共责任，抵制商业资本过度侵蚀

UNESCO 强调政府在 AI 教育治理中的核心公共责任，主张政府主导 AI 教育治理，抵制商业资本对教育公共性的过度侵蚀。乔治·西门子在“人与机器：新兴人工智能能力的政策影响”一文中指出：“人工智能的快速发展，涉及伦理考量、数据隐私和数字鸿沟等方面的问题。政府必须营造鼓励创新的环境，并确保所有学生，无论其社会经济背景或地理位置如何，都能公平地获取基于人工智能的学习工具和资源。”与此同时，当前商业资本在 AI 教育领域过度扩张，将教育视为盈利市场，主导 AI 技术的研发与应用，挤压公共教育空间，损害教育公共利益。为此，UNESCO 主张，政府必须承担 AI 教育的公共责任，加大公共投入，主导公共 AI 教育系统研发，规范商业企业在教育领域的行为，禁止商业资本垄断 AI 教育资源、滥用学生数据、侵蚀教育公共性，确保 AI 教育始终服务于公共利益而非商业利益。

（三）本土适配、文化尊重，反对全球标准化治理

UNESCO 反对 AI 教育的全球标准化治理模式，主张本土适配、文化尊重、多元治理，尊重不同国家、不同文化的教育特色与治理需求。UNESCO 在该文集中强调，全球各国的教育体系、文化传统、社会语境差异巨大，单一的 AI 教育治理模式无法适配所有国家，必须推动 AI 教育治理的本土化、语境化适配。各国应结合本土教育实际、文化传统、发展需求，制定符合本国国情的 AI 教育治理政策，同时尊重本土知识、小众语言与传统教育智慧，避免西方主导的标准化治理模式消解全球人才培养的多样性，培养符合本土发展需求的人才。

六、启示与建议

近年来，我国加强与 UNESCO 等多边机构的合作，深度参与国际教育规则制定，提供中国的教育发展标准和方案，受到了国际组织的深入认可。近期，我国出台了《国务院关于深入实施“人工智能+”行动的意见》《关于加快推进教育数字化的意见》《“人工智能+教育”行动计划》等一系列政策，推动人工智能在教学、学习、评价、治理中的全场景应用。UNESCO 这部文集集中的基本立场与主张，也与我国的 AI 教育理念高度契合。该文集所汇聚的关于 AI 人才培养的国际前沿观点，足以为我国人工智能时代的人才培养提供相应的参考。有鉴于此，结合我国 AI 人才培养发展，提出如下建议。

（一）坚守 AI 时代人才培养的公共性与人文性，抵制技术与商业异化

作为社会主义国家，我国的人才培养应始终坚守公共性本质与人文性内核，将其作为 AI 时代教育的根本遵循。首先是强化教育公共性导向，明确 AI 教育是公共教育服务的重要组成部分，加大政府公共投入，主导公共 AI 教育平台、教育专用大模型的研发与应用，严格规范商业企业进入教育领域的行为，禁止商业资本垄断 AI 教育资源、滥用学生数据牟利，维护教育的公共属性。其次是坚守人文性底线，将马克思主义关于人的全面发展理念、教育家精神等置于 AI 教育核心，反对用 AI 替代教师、用算法消解教育温度，推动 AI 技术服务于立德树人根本任务，实现技术理性与人文关怀的统一。在拥抱 AI 技术的同时警惕技术至上主义，摒弃

“技术决定论”的错误理念，确保 AI 技术始终服务于教育规律与育人目标，避免技术绑架教育。

（二）以公平包容为核心，缩小 AI 时代人才培养的群体差距

将公平包容作为 AI 教育发展的首要目标，着力缩小城乡、区域、群体之间的 AI 教育差距。一方面，优先保障农村与欠发达地区权益，加大对农村、边远、民族地区的 AI 教育基础设施投入，确保欠发达地区学生平等享有 AI 学习资源。另一方面，关注弱势群体公平，保障残障学生、贫困学生、留守儿童的 AI 教育权益，开发适配特殊需求的 AI 学习工具，将 AI 素养教育纳入全学段课程，确保所有学生具备基础 AI 能力。

（三）强化师生主体地位，夯实 AI 时代人才培养的核心支撑

有鉴于 UNESCO 关于 AI 时代教学与学习的剖析，在推进 AI 教育教学过程中，应始终坚守教师与学生的主体性。明确教师是 AI 教学的主导者，AI 是辅助工具，严禁用 AI 虚拟教师替代人类教师，加强教师 AI 素养与伦理能力培训，将 AI 教学能力纳入教师专业发展体系。对于学生学习而言，应抵制认知卸载，引导学生合理使用 AI 工具，培养批判思维、创新能力与自主学习能力，将 AI 作为学习助手而非替代者，强化学生的学习主体性，培养学生的 AI 核心素养，为未来的人机共生奠定能力基础。

（四）创新人才培养评价体系，发挥评价的育人导向作用

针对 AI 对传统评价体系的冲击，借鉴 UNESCO 主张，加快推进我国人才评价体系革新。一是丰富评价体系，加强表现性评价、项目式评价、过程性评价等评价方式，真实反映学生的素养水平。二是构建人机协同评价模式，发挥 AI 在过程性数据收集、客观题批改中的优势，将主观评价、素养评价、人文评价交由教师主导，坚守教师的评价主导权。三是强化评价公平性，对 AI 评价系统进行算法偏见审计、公平性审查，消除评价中的性别、地域、阶层偏见，保障所有学生的评价公平权益，引导人才培养向素养导向转型。

（五）构建多元共治的人才培养治理体系

我国应完善 AI 人才培养治理体系，借鉴 UNESCO 的治理主张，构建多元共治、伦理先行、法治保障的治理格局。一是完善多元主体参与机制，形成政府统筹、教育部门主导、学校落实、教师学生参与、企业规范、社会监督的共治体系，保障一线教育者与学生的话语权。二是强化法治监管，完善 AI 教育相关法律法规，明确数据隐私、责任划分、商业行为边界，严厉打击数据滥用、算法歧视、商业侵蚀等行为。三是推动本土适配与国际合作，结合我国教育国情制定治理政策，同时加强与 UNESCO 等国际组织的合作，参与全球 AI 教育治理，分享中国经验，借鉴国际共识，推动构建人类命运共同体导向的人才培养新生态。新的历史时期，我国作为全球 AI 教育发展的重要参与者与引领者，宜充分吸收 UNESCO 关于 AI 教育的核心主张，立足我国教育国情，坚守育人为本、智能向善的理念，平衡技术创新与教育本质、效率提升与公平包容、商业活力与公共利益，推动人工智能与教育深度融合。让 AI 技术真正服务于教育公平、质量提升与人的全面发展，为全球 AI 时代的人才培养贡献中国智慧与中国方案。

作者信息

王小飞 中国教育科学研究院比较教育研究所所长、研究员

赵章靖 中国教育科学研究院比较教育研究所副研究员

李佐文

欧美国家人工智能治理话语的计算分析： 主题特征与演化趋势

来源 | 《教育国际交流》2026 年 03 期



北京外国语大学教育学院党总支书记、人工智能与人类语言重点实验室主任 李佐文

1. 引言

人工智能技术的快速发展，特别是大语言模型时代的到来，对经济社会发展和人类文明进步产生深远影响，给世界带来巨大机遇。与此同时，人工智能技术也带来难以预知的各种风险和复杂挑战。涌现出来的隐私伦理、算法伦理、技术伦理等问题对个人与社会带来的潜在不可控风险也引起全世界的共同关注。人工智能治理关乎全人类命运，是世界各国面临的共同课题。根据斯坦福大学的统计，推出含有人工智能内容立法的国家已由 2022 年的 25 个增加至 2023 年的 127 个。经济合作发展组织、联合国教科文组织、欧洲委员会以及七国集团、二十国集团等也在多边层面出台了人工智能治理的原则与框架。目前全球人工智能治理相比于人工智能技术的迅速发展还具有滞后性。国际社会虽然已就治理人工智能达成

共识，但风险管理和监管仍是行业内的薄弱环节。

欧盟和美国是人工智能技术领域的领先者，在全球人工智能治理中扮演重要角色。美国人工智能相关法规的数量在过去五年中显著增加。2023 年人工智能相关法规数量达到 25 个，而 2016 年仅有 1 个。仅 2023 年一年人工智能相关法规总数就增长了 56.3%，显示出美国政府对人工智能治理的高度重视和快速响应。2024 年 5 月，全球首部综合性人工智能监管法案《人工智能法案》（AI Act）正式生效，为全球范围内的人工智能监管提供了重要参考和借鉴，标志着人工智能治理成为全球治理的全新领域。对比并学习国外人工智能政策规划的先进经验，是中国进一步发展人工智能产业的必然要求。

人工智能政策文本是研究人工智能治理话语的重要素材。话语的主题计算在话语整体层面上进行，是对概念意义的抽象和概括。主题意义是从整合话语的各个组成部分中提取出来的能够涵盖整个语篇内容和主旨。计算人工智能政策文本的主题特征，挖掘分析其中的主题意义，对洞察预见人工智能治理话语的脉络和趋势具有重要意义，对中国人工智能治理策略和发展战略的制定具有重要参考价值。在现今国际新秩序构建渐趋深化、国际权力关系调整变化的形势下，争取人工智能领域应有的国际话语权并加强人工智能话语权建设，不仅可以有效推动我国科技水平的持续进步，还能够显著增强我国的国际影响力。

人工智能的发展与监管是国家兼而有之的两个核心政策目标。在发展和监管的双重政策目标框架下，各国政府会形成各有侧重的发展政策与监管政策。国家会基于对技术创新路径和技术风险认知的理解，确立发展和监管人工智能的整体政策目标和具体政策手段。欧盟人工智能治理的突出特征之一是“伦理优先”，即构建完善的伦理原则及法律体系；美国则坚持“创新优先导向”，尽力维护国家人工智能创新能力。

本文围绕欧盟和美国的重要人工智能政策文本进行合并研究，通过动态主题模型探析欧盟和美国人工智能治理话语的主题特征，厘清欧盟和美国在人工智能治理领域共同关注的主题。并通过分析主题演化趋势，展望人工智能治理话语的发展前景。

2. 研究设计

2.1 研究方法思路

为分析欧美人工智能治理话语的主题及其演进过程，本研究采用动态主题模型（Dynamic TopicModel，简称DTM），主要包括主题识别、主题演化分析两项主要内容，从静态和动态两个方面分析欧美人工智能政策主题的结构、特征和发展趋势。LDA模型是一种非监督机器学习的文本挖掘方法，基于狄利克雷分布和贝叶斯算法提取文本包含的潜在主题信息和主题的词汇分布。DTM模型由Blei等在LDA模型的基础上提出，属于LDA模型的变体，通过引入时间因素，以时序为依据对文档集进行切片，并通过比较相邻时间文本主题分布差异模拟主题的演化规律。

本研究首先对数据进行预处理，包括数据清洗、分词、去除停用词等。随后利用DTM进行主题建模，通过评估不同主题数下模型的主题一致性参数并结合实际演化效果确定最佳主题数。在此基础上，依据“主题—词项”和“文档—主题”的概率分布提取各主题及高频率词汇。最后，结合时间维度计算主题强度，从而揭示各主题的客观演变规律。

2.2 数据采集

本研究的政策文本选取自欧盟和美国在2016年1月至2024年8月间出台的代表性人工智能治理政策。为确保选取的政策文本具有研究价值，采取以下3个筛选原则：一是政策文本标题或内容必须明确与人工智能治理相关。二是具有重要性，政策文本均来自政府官网发布的公开文件，并且对全球或本国人工智能治理的后续政策制定或立法产生深远影响。三是具有权威性，政策文本必须是由政府、议会等官方机构发布。基于以上原则并结合相关研究，在欧盟和美国的政策法规官方门户网站上共查找并下载24份相关政策文件，见表1和表2。

表₁ 2016—2024 年欧盟发布的人工智能治理重要政策文本

序号	名称	时间
1	<i>Draft Report with Recommendations to the Commission on Civil Law Rules on Robotics</i>	2016 年 5 月
2	<i>Artificial Intelligence for Europe</i>	2018 年 4 月
3	<i>Coordinated Plan on Artificial Intelligence</i>	2018 年 12 月
4	<i>Ethics Guidelines for Trustworthy AI</i>	2019 年 4 月
5	<i>White Paper on Artificial Intelligence</i>	2020 年 2 月
6	<i>Proposal for a Regulation Laying Down Harmonised Rules on Artificial Intelligence</i>	2021 年 4 月
7	<i>Directive of the European Parliament and of the Council on Adapting Non-contractual Civil Liability Rules to Artificial Intelligence</i>	2022 年 9 月
8	<i>AI Liability Directive</i>	2022 年 9 月
9	<i>Proposal for a Regulation of the European Parliament and of the Council Laying down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act) and Amending Certain Union Legislative Acts — General Approach</i>	2023 年 6 月
10	<i>ENISA Threat Landscape 2023</i>	2023 年 10 月
11	<i>Artificial Intelligence Act</i>	2024 年 5 月

表₂ 2016—2024 年美国发布的人工智能治理重要政策文本

序号	名称	时间
1	<i>Preparing for the Future of Artificial Intelligence</i>	2016 年 10 月
2	<i>The National Artificial Intelligence Research and Development Strategic Plan</i>	2016 年 10 月
3	<i>Artificial Intelligence, Automation, and the Economy</i>	2016 年 12 月
4	<i>FUTURE of Artificial Intelligence Act of 2017</i>	2017 年 12 月
5	<i>Artificial Intelligence Emerging Opportunities, Challenges, and Implications for Policy and Research</i>	2018 年 6 月
6	<i>Maintaining American Leadership in Artificial Intelligence</i>	2019 年 2 月
7	<i>A Bill to Establish the National Artificial Intelligence Initiative, and for Other Purposes.</i>	2019 年 11 月
8	<i>Memorandum for the Heads of Executive Departments and Agencies Guidance for Regulation of Artificial Intelligence Applications</i>	2020 年 11 月
9	<i>Artificial Intelligence Training for the Acquisition Workforce Act</i>	2021 年 12 月
10	<i>Blueprint for an AI Bill of Rights</i>	2022 年 10 月
11	<i>Artificial Intelligence Risk Management Framework</i>	2023 年 1 月
12	<i>National Artificial Intelligence Research and Development Strategic Plan 2023 Update</i>	2023 年 5 月
13	<i>Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence</i>	2023 年 10 月

2.3 数据预处理

本研究通过 Python 进行 DTM 主题建模，按照以下 3 个步骤进行预处理。

首先是清洗语料。文本数据中包含部分与主题无关且会对模型效率产生影响的噪声数据，如标点符号、数字和无实际意义词汇等。去除完成再使用 NLTK 库进行分词，将文本拆分为单独的英文词汇。分词后的文本更易于进行词频统计、特征提取等基础操作，能够更好地反映语义信息。获取分词词表后再对停用词进行过滤，停用词是指在文本中频繁出现但语义信息含量较低的词。引入停用词表

以增强文本的语义显著性，提高主题模型的准确性和效率。最后结合 NLTK 库中的 WordNetLemmatizer 实现词形还原，将各种形式的英文词汇还原为一般形式。

其次是统计词频。主题建模的关键在于确定“主题—词项”与“文档—主题”的概率分布。文本中出现频率过低或无实际语义贡献但出现频率过高的词语，均会显著影响模型的计算效率和计算结果。本研究对预处理后词语的词频进行统计，在主题建模前删除月份（如 April）等低频词和国家名称（如 America）、组织名词（如 EU）等无意义高频词。

最后是构建词典，将文本表示为向量。文本向量化指将文本数据转换为数值型数据的过程。本研究使用 Gensim 库中的 Corpora 模块，创建用于主题分析的词典和语料库，构建词袋模型（Bagof Words）将分词后的文本表示为词频向量，以便进行下一步的计算分析。

2.4 最优主题数确定

本研究使用 Gensim 和 pyLDAvis 进行主题建模和可视化分析，使用 Gensim 库中的子模块 Models 中的主要模型 LdaSeqModel 进行 DTM 主题建模，通过计算一致性确定最优主题数，主题一致性计算结果见图 1。

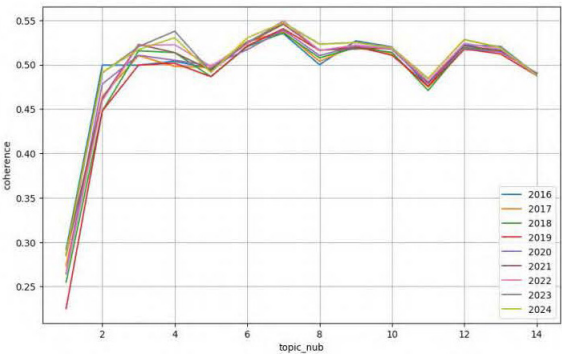


图1 主题一致性参数

表3 “主题—词项”分布表

编号	主题内容	主题高频词(节选)
Topic0	Ethical Principles and Trustworthy Development Framework for AI (人工智能伦理与可信赖人工智能发展框架)	system, use, data, human, ethical, development, ensure, trustworthy, assessment, public
Topic1	Socioeconomic Impacts of AI and Safety Considerations (人工智能的社会经济影响与安全考量)	market, law, economic, harm, impact, liable, safety, risk, consumer, uncertainty
Topic2	AI Liability Distribution and Burden of Proof Challenges (人工智能责任分配与举证难题)	burden, proof, relevant, use, law, certain, specific, risk, particular, harm
Topic3	Cybersecurity Threats and Responses in AI (人工智能网络安全威胁与应对)	information, threat, data, use, landscape, security, management, cybersecurity, ransomeware, attack

Topic4	AI Regulation and Conformity (人工智能的监管与合规性)	system, regulation, high-risk, conformity, set, appropriate, competent, authority, ensure, provider
Topic5	AI Risk Assessment and Economic Compensation Mechanisms (人工智能风险评估及经济补偿机制)	compensation, economic, liable, victim, approach, initiative, civil, harmonised, trust, directive
Topic6	AI Research and Future Development Trends (人工智能技术研发与未来趋势)	research, technology, data, science, government, future, automation, energy, system, human

表4 主题强度

编号	主题	强度
Topic0	Ethical Principles and Trustworthy Development Framework for AI (人工智能伦理与可信赖人工智能发展框架)	4.897727124788489
Topic1	Socioeconomic Impacts of AI and Safety Considerations (人工智能的社会经济影响与安全考量)	0.6383945269582859
Topic2	AI Liability Distribution and Burden of Proof Challenges (人工智能责任分配与举证难题)	0.610946128459793
Topic3	Cybersecurity Threats and Responses in AI (人工智能网络安全威胁与应对)	1.4413065245195824
Topic4	AI Regulation and Conformity (人工智能的监管与合规性)	3.002729167158188
Topic5	AI Risk Assessment and Economic Compensation Mechanisms (人工智能风险评估及经济补偿机制)	0.8687314711136163
Topic6	AI Research and Future Development Trends (人工智能技术研发与未来趋势)	9.540165057002046

一致性 (Coherence) 是衡量主题模型水平的一个指标, 可以通过计算某个主题中前几个关键词的共现关系来度量, 可以反映模型生成的各个主题词之间的语义相关性。最佳的主题数量通常是一致性最高的主题数。一个模型的一致性参数越高, 生成的主题在语义上越紧密, 即模型质量好。如图 1 所示, 当主题数为 7 时主题一致性参数最高, 因而本研究设定主题数为 7。

3. 主题识别

通过 DTM 建模, 对文本数据进行主题分析。根据计算结果, 提取 7 个主题以及每个主题的特征词分布, 节选其中的高频率词项根据出现概率由高到低排列。根据每个主题下的高频词, 总结和归纳出具体主题内容, 见表 3。

主题强度的大小就是主题在某一个时间窗口上受到的关注程度, 即在某一时间窗口上某主题强度与包含该主题的文档数目成正比, 主题强度越大越有可能被认为是热点主题。各主题对应的强度如表 4 所示。从主题强度分布来看, 人工智能治理话语的核心关注点集中在人工智能伦理与可信赖人工智能发展框架和人工智能技术研发与未来趋势 (Topic0 和 Topic6)。人工智能的监管与合规性 (Topic4) 和

人工智能网络安全威胁与应对 (Topic3) 受到一定重视, 人工智能责任分配与举证难题 (Topic2)、人工智能风险评估及经济补偿机制 (Topic5) 和人工智能的社会经济影响与安全考量 (Topic1) 领域也是主要议题。

从上述数据可知, 欧美人工智能治理话语内含技术发展与社会治理并重的趋势。话语主题主要集中在伦理与监管框架、社会经济影响、责任划分、网络安全、监管合规、风险评估与赔偿机制以及技术研发等多个领域。从伦理框架的构建到网络安全的应对, 从责任划分到经济赔偿机制的完善, 人工智能治理话语的内容日益全面化。

Topic0 为人工智能伦理与可信赖人工智能发展框架, 其高强度反映了这一领域的重要性。Ethical 和 trustworthy 是欧美人工智能治理话语中的高频词。人工智能技术的应用引发大量伦理争议, 人工智能伦理问题逐渐成为全球性讨论的热点。例如, 人工智能系统在数据训练中的偏见问题导致社会公平性的质疑。面向使用者的人工智能服务可能在无意间强化性别、种族或经济阶层的不平等。围绕 Trustworthy AI 制定相关的框架和原则是欧美国家关注的核心议题之一。2019 年欧盟率先提出《可信赖的人工智能伦理准则》(Ethics Guidelines for Trustworthy AI), 强调人工智能系统应具有透明性、公开性和可问责性。2020 年后, 欧盟在《人工智能白皮书》(White Paper on Artificial Intelligence) 和《人工智能法案》(AI Act) 中倡导建立卓越生态系统和人工智能监管框架, 精细划分人工智能风险等级, 并制定针对性的监管措施, 至此实现全流程的伦理治理。美国分别于 2022 年和 2023 年发布的《人工智能权利法案蓝图》(Blueprint for an AI Bill of Rights) 和《关于安全、可靠、可信地开发和使用人工智能的行政命令》(Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence) 都讨论了可信赖人工智能的相关框架和原则, 展望了可信赖人工智能的理想状态。

Topic1 为人工智能的社会经济影响与安全考量。通过高频词 market, economic, harm, impact, safety, risk 和 uncertainty 等, 可以看出人工智能技术在改变社会和经济格局方面表现出显著的双刃剑效应。一方面, 人工智能提升了生产效率和推动了新兴产业的发展; 另一方面, 其广泛应用也引发劳动力市场的

结构性变化以及市场的不确定性问题。随着人工智能技术在各领域的深入应用，人工智能对劳动力市场的潜在影响逐渐成为社会讨论的焦点。美国出台的《人工智能劳动力培训法案》（*Artificial Intelligence Training for the Acquisition Workforce Act*）关注如何通过教育和培训提升公众的人工智能素养，并采取措施减少技术进步带来的不平等问题，减弱人工智能对劳动力市场的冲击。这一主题强度相对较低，表明尽管这一领域涉及广泛，但欧盟和美国更聚焦于技术发展，社会、经济层面的具体措施有待深化。

Topic2（人工智能责任分配与举证难题）、Topic5（人工智能风险评估及经济补偿机制）两个主题强度相对较低。两个主题中的高频词 *burden*, *proof*, *regulation*, *high-risk*, *compensation* 和 *victim* 等共同体现了人工智能技术的发展速度远远超过现有法律政策监管框架的更新速度。尤其是当人工智能出现决策失误或导致损害时，责任认定与问责机制尚不完善，受害方的权益难以保障。人工智能责任分配与举证难题这一领域虽然重要，但在政策制定中尚未成为焦点的原因与人工智能技术的复杂性有关，算法的“黑箱性”使得责任分配极其困难。

Topic3（人工智能网络安全威胁与应对）主题强度适中，说明欧美人工智能治理话语在这一主题上具有一定关注度。相关高频词有 *data*, *ransomware* 和 *attack* 等。人工智能系统在大规模数据处理中展现出强大算力，但深度伪造技术、勒索软件、数据泄露、人工智能技术驱动的网络攻击等问题频发，反映了人工智能技术浪潮下网络安全面临新的挑战。

Topic4（人工智能的监管与合规性）强度突出，说明监管发展与合规要求的显著重要性。高频词如 *system*, *regulation*, *high-risk*, *conformity*, *authority* 和 *ensure* 等表明欧美人工智能治理话语共同关注人工智能监管的必要性。人工智能将越来越普遍地运用于人类生产和生活中，比如自动驾驶汽车、人脸识别等，随时可能危害人类社会的安全和侵犯人类基本权利。但由于其技术原理的复杂性、模糊性和不可解释性，现有法律规则供给不足，监管机构面临规制困境。监管与合规性作为确保人工智能技术安全、可靠和公平应用的基石，其显著重要性不容忽视。尽管欧盟和美国在人工智能的监管与合规性方面采取了不同的策略和法律

框架，但在核心理念和目标上存在共性。两者都强调通过系统性的监管、识别和管理高风险领域、建立合规机制、明确管辖权确保人工智能技术的应用符合社会利益，推动人工智能向善发展。

Topic6 是人工智能技术研发与未来趋势，也是欧美人工智能治理话语的核心内容。通过高频词研究 research, technology, science, government, automation, energy 和 system 等，可以看出欧美政府在人工智能技术研发和未来人工智能的发展中仍会起到重要作用以维持技术竞争力。例如，美国在 2016 年和 2023 年发布的两版《国家人工智能研究和开发战略计划》（National Artificial Intelligence Research and Development Strategic Plan）中均明确了以基础研究为核心的人工智能发展战略。另外，人工智能技术研发的未来趋势关注能源效率提升。人工智能技术的运行和数据处理依赖于高能耗计算设施，加快发展绿色人工智能技术，减少能源消耗是未来人工智能研发的一大重点。

4. 主题演化分析

主题演化数据的结果表明，欧美人工智能治理话语的主题自 2016 年以来经历了显著变化，不同主题的关注度随着时间推移呈现出周期性的波动。人工智能技术的突破发展和应用扩展更深刻地推进了全球人工智能治理政策的发展，塑造了人工智能治理话语的演进，丰富了人工智能治理话语体系。

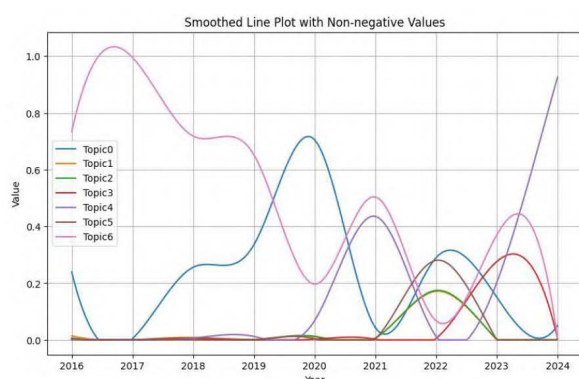
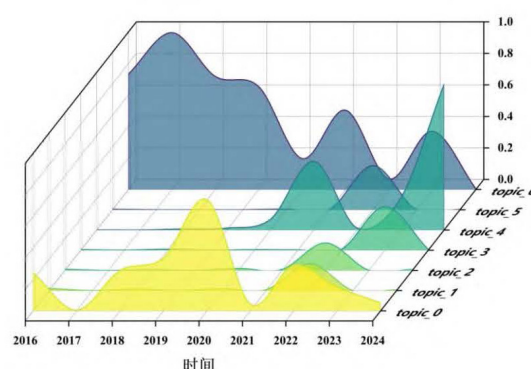
表。 研究主题强度演化

时间	Topic0	Topic1	Topic2	Topic3	Topic4	Topic5	Topic6
2016	0.239553	0.01416	0.006463	0.003400	0.000603	0.002998	0.732823
2017	0.00417	0.000039	0.000039	0.000039	0.000039	0.000039	0.995636
2018	0.255836	0.008319	0.005185	0.006014	0.006397	0.000089	0.71816
2019	0.334888	0.000009	0.000009	0.000309	0.012247	0.000009	0.65253
2020	0.705034	0.011129	0.012566	0.00261	0.068063	0.003977	0.196621
2021	0.043644	0.006653	0.006344	0.001801	0.435334	0.002587	0.503637
2022	0.289585	0.171002	0.174111	0.005118	0.012532	0.280397	0.067255
2023	0.150206	0.000671	0.000897	0.275943	0.202117	0.000171	0.369996
2024	0.050699	0.004834	0.004863	0.005734	0.926112	0.001258	0.00650

人工智能技术的发展经历了多次技术突破和里程碑事件，尤其是 2020 年 7 月 GPT-3 模型的发布以及 2022 年 11 月 ChatGPT 的问世从技术层面展示了语言智能的卓越能力（李佐文 2022）。GPT-3 的发布标志着大模型技术初步成熟，而 ChatGPT 则以其强大的生成能力带来里程碑式的技术进步。由图 2 可知，这两个时间节点也成为欧美政策制定者重新审视人工智能风险与机遇的重要契机。

首先，人工智能技术研发与未来趋势（Topic6）始终是欧美人工智能治理话语的核心内容，其主题强度虽有波动，但 GPT-3 和 ChatGPT 的发布掀起生成式人工智能的大规模研发和创新浪潮。其次，伦理与社会影响议题有一定升温。人工智能伦理与可信赖人工智能发展框架（Topic0）和人工智能的监管与合规性（Topic4）的主题强度增加，成为欧美人工智能治理话语讨论的重点之一。但 2022 年 11 月 ChatGPT 面世后，人工智能伦理与可信赖人工智能发展框架（Topic0）的主题强度下降。可以看出伦理问题尽管仍然重要，但其关注度被人工智能技术广泛应用及其带来风险所“稀释”。生成式人工智能的大规模应用直接影响劳动力市场、教育体系和社会分工，关于人工智能的社会经济影响与安全考量（Topic1）的讨论从 2020 年后开始显著上升。责任分配与举证难题（Topic2）在 2021 年逐步成为关注点，2022 年后继续升温。人工智能风险评估及经济补偿机制（Topic5）在人工智能技术逐渐成熟后成为新兴议题。到 2023 年，该主题的热度显著提升。人工智能大规模应用的背景下如何合理进行风险评估和补偿机制的制定引起热议。人工智能网络安全威胁与应对（Topic3）在人工智能应用中始终是重要议题，ChatGPT 生成内容的高度复杂性和不可预测性使得此类担忧进一步放大。人工智能的监管与合规性（Topic4）的主题强度在 2022 年后迅速上升，人工智能的不可控风险迫切呼监管治理体系的构建和完善。

根据主题强度演化趋势图（见图 2），预计未来，人工智能的社会经济影响与安全考量（Topic1）和人工智能的监管与合规性（Topic4）将继续占据主导地位。目前针对人工智能显性风险的政策条目逐渐增多，隐性风险规制与监管政策较为欠缺。人工智能伦理与可信赖人工智能发展框架（Topic0）的主题强度可能有下降趋势，但这并不意味着伦理问题的重要性减弱，表明此议题可能将逐步更加

图₂ 主题强度演化曲线图图₃ 主题演化瀑布图

细化或整合到其他主题（如监管与合规性、责任分配）中。人工智能责任分配与举证难题（Topic2）和人工智能网络安全威胁与应对（Topic3）的讨论度将逐步上升，并与技术发展和政策监管形成深度关联。人工智能风险评估及经济补偿机制（Topic5）作为补充性议题在高风险应用场景中尤为关键，未来将得到更多关注。人工智能技术研发与未来趋势（Topic6）将长期保持热度，基础技术研究是人工智能行业实现可持续发展的关键与动力。人工智能技术的未来发展依赖于技术与治理的平衡推进。监管体系全球化与区域协作、具体化的伦理与责任机制、社会经济适应性政策、安全与风险治理创新等要素将共同确保智能向善。

5. 结束语

人工智能已经成为关乎国家安全和经济发展的关键技术，人工智能治理也已经成为全球治理的重要领域。本文以欧盟和美国人工智能治理的标志性政策文本为研究对象，通过 DTM 模型对欧美人工智能治理话语进行主题识别与主题演化分析。

结果发现，人工智能技术研发、人工智能风险管理、人工智能监管与合规要求以及人工智能伦理等是欧美人工智能治理话语共同关注的议题。未来，欧美人工智能治理话语的演进将聚焦在统筹安全与发展，平衡监管与创新。通过监管体系全球化与区域协作、伦理与责任机制具体化、社会经济适应性政策、安全与风险治理创新等举措，构建全面细化的可信赖人工智能治理框架，确保人工智能技术的健康发展。

中国作为人工智能强国，通过吸收欧美经验并结合自身实践，可探索出一条既保障安全与伦理，又促进技术发展的平衡治理道路，构建发展与安全并重的人工智能治理体系。加强人工智能的安全风险评估和管理，致力于构建国际人工智能安全的通用标准。在此基础上，以更加积极的姿态参与人工智能的全球治理，为推进人工智能的全球治理继续提供中国方案，贡献中国智慧，提升中国在人工智能治理领域的国际话语权。

作者信息

李佐文 北京外国语大学教育学院党总支书记、人工智能与人类语言重点实验室主任

张一凡 北京外国语大学中国外语与教育研究中心



龙宝新

生成式人工智能时代教育创新范式的演进路向

来源 | 《中国教育科学（中英文）》2026 年 03 期



陕西师范大学教师教育协同创新中心主任，教育学部教授、博士生导师 龙宝新

创新是推进教育实践系统性重构的内在动能，也是教育持续迭代升级的隐秘力量，教育和创新共同决定着未来教育的样态与走向。缺乏创新意愿的教育是灵魂缺位的教育。缺乏创新举措的教育是没有未来的教育。教育创新不仅彰显教育本质的能动力量，还是教育事业永葆生命活力的内生动能。当前，教育正步入由生成式人工智能（generative artificial intelligence，简称 GAI）主宰的新技术时代，教育创新新形态悄然诞生。如何因应时代特征和国际形势，推进教育创新范式演进，是一项事关我国教育事业发展的重要实践。

一、教育创新范式回眸

范式是“用来观察和理解世界的一套完整的概念框架”或思想体系，范式的形成标志着人们对事物的认识已步入一个相对稳定的历史水平，核心创新元素正

是一种教育创新范式的硬核构成。一部教育发展史就是一部教育改革史、一部教育创新史，其中质变飞跃的节点将教育创新实践划分为一系列连续性发展阶段，每个阶段都被一种核心教育创新范式所主宰。在不同历史时期曾经出现过三种典型的知识生产方式，即经验形态的知识生产、原理形态的知识生产、信息技术支撑的交叠形态知识生产。与之相对应，教育创新范式大致包括三类，即经验式教育创新、反思式教育创新、知识集群式教育创新。教育创新凭依经验量变积累、知识媒介驱动、知识创新体搭载等创新元素，依次更迭创新范式。

（一）经验式教育创新

创新是古今中外普遍存在的社会变革现实，是人类发展的各个阶段都在发生的普遍现象。创新活动的发展呈现日益专门化、深刻化，日益融入日常生活的趋势。教育创新是教育生活的永恒主旋律。在科学知识诞生之前，普遍存在于教育实践领域之中的是一系列教育见识、教育感验、教育随想，这就是教育经验。借助对教育活动、教育事象的感性看法、直觉反映、表象表征来建构教育世界，正是原始教育实践的运行方式，也是原生态教育创新的基本形态。在漫长的古代社会，教育创新凭借的是教育经验的自然积累，没有这种原生、原始、原本的教育经验创新形态，教育世界根本难以生存延绵。古代教育实践者还有更多教育创新工具，如象征、表象、意象、隐喻、类比、知觉、直觉等知识替代性思维工具，它们构成了教育实践者的非知识类教育创新手段集合，从教育实践最底层支撑着教育世界的演进与更迭。经验式教育创新范式有三个共同特征。一是知行一体性，即触发教育创新的工具手段，如象征、表象、意象等内生于教育实践母体之中，与教育实践情景融为一体，“干中思”“干中悟”“干中创”是其典型存在样态，体现出“见识行动”一体化统合的特征。二是生长连续性，这种教育创新不同于知识媒介的教育创新形态，教育实践者无须抽身事外便可进行，教育创新实践始终保持着连续性、不间断性的特征。三是情境栖身性，原生的经验式教育创新始终栖身于教育情境，并被教育生活与教育实践场景所催生、所承载、所滋养，情境性教育智慧是教育创新实践的中流砥柱。

（二）反思式教育创新

随着以哥白尼日心说为代表的第一次科学革命发生，经验哲学宣告寿终正寝。

德国古典哲学展开了“思维与存在关系”的论辩，标志着科学知识时代的正式来临。在这一背景下，基于知识生产媒介的教育创新方式成为教育革新的新引擎，反思代替反映、概念代替见识、意识代替感觉、建构代替回应成为教育创新的新范式。教育创新是对原有不合理的理论观点、思想方法、技术手段进行突破和超越，以便在学校变革与发展中更有效地实现教育目的的一种创造性活动，其实质在于不断地反思批判和重新选择。教育反思的功能是加速教育事物、教育实践自然存在样态的裂变，将原本完整的教育经验、教育实践一分为二。教育创新必须时刻面对“理论实践”“知识行动”间的紧张与张力关系，二者若即若离的矛盾与焦虑关系推动了教育创新实践的持续发生。

教育知识创新是教育实践创新的先导，教育知识的创造力成为决定教育实践创造水平的关键性参量。在科学知识诞生后，教育创新图景发生了质的改观：教育创新实践与教育知识生产之间构成同命运、同呼吸、同生产的关系。改造教育知识的生产方式、提升教育知识的产能，成为教育创新实践的全新支撑点。符号化、概念化、抽象化教育知识的诞生并没有完全消除经验式教育创新，反思式教育创新具有教育知识生产的媒介性和教育创新实践的分裂性两种特性，每一个教育创新行动意图的实现都可通过教育知识界与教育实践界的跨界转生、双向互动来推动。

（三）集群式教育创新

从20世纪70年代开始，人类社会进入后工业社会，知识、信息、数据成为重要的生产资源与生产力要素，知识社会化成为这一新时代的根本特征。在这一背景下，知识生产社会化、跨界化、协同化成为教育创新的新特点，一种以集群组织、创新网络、分形研究为代表的集群式教育创新模式应运而生。在后工业社会，教育创新组织变革最为抢眼，“UG—S”（大学政府中小学）协作体、产学研政一体化、教育知识创新共同体、跨界教育创新联盟等新型组织模式代替了单一的教师创新者或教育研究者，一跃成为教育创新的主要形态。集群式教育创新活动横跨多个教育产业行业部门，教育知识的生产、应用、验证、再创新向着一体化方向大步迈进，教育知识创新组织空前开放，“从只重知识的生产与传播转向同时注重知识的生产、传播和应用，消除知识生产与知识应用之间的隔阂”，在企业与大学之间流动的教育科技人员发挥了“助产士”的作用。“大学政府产业公众”教育知识生产集群成形，标志着教育创新步入多主体跨界协同的新阶段。

教育创新实践发展呈现出高度分化与高度整合的历史性特征，数据、信息、知识、智慧等多层级的信息资源成为串联教育创新实践的连心锁与连通器，数据链、信息流、知识连续体将各界教育创新实践贯通一体。在多主体、网络状、集群态的教育创新组织中，几乎每一个教育实践参与者都可能成为一个创新主体，每一个具体教育实践链环上都遍布教育创新的要素、点位，所有教育创新实践借助信息网、数据链、云空间被串联耦合起来，进而实现了教育创新要素的最优化耦合状态。正是有了公共教育创新网络空间的加持，跨界教育实践主体间的思想联通、智慧联网、观点共享变得更加便利，教育创新更容易在协同攻关点位上集中爆发，进而催生出高创造性、高价值性、高维度性的教育创新成果。在后工业社会，随着生成式人工智能的出现，教育创新者作为教育创客，不仅可以借助创新主体间联网发挥集群创新的功能优势，还可以借助智能机器人、智能学伴助理提升自身的教育创新能力，教育创新实践随之实现了在“横向联合”与“纵向深化”两个维度上的质变与扩容。在这一阶段，“教育创新应该像血液一样成为所有教育行为者和教育机构的功能，流经每一个体与组织细胞，激发教育的全面活力”，成为教育实践的关键性构成单元。集群式教育创新是迄今为止人类历史上最高级的一种教育创新形态，代表着当前教育创新范式演化的最高水平。

二、生成式人工智能：一个全新的教育创新元

在集群式教育创新阶段，教育创新真正步入多向化、多元化、多极化发展阶段，每一个新创元素的加入都可能改变教育创新的模态、架构、组织与效能。有学者指出，以 ChatGPT 为代表的生成式人工智能以其自主学习能力、自注意力、推理能力等突破性智能技术特征引发了全球性的“AIGC+”时代浪潮。生成式人工智能的加入，将彻底改变未来教育活动的组织架构。创新意味着组合之前未有过的生产条件与生产要素，使之形成一种新的生产函数并进入已有的生产体系。生成式人工智能在教育领域的广泛应用，意味着一种全新教育创新组合体的出现，将全方位影响教育创新方式，带来教育生态的演化提速。

（一）升级教育创新方式

生成式人工智能重塑教育创新的方式较为独特：它不是以完全独立创新主体的身份参与教育实践创新活动，而是依存于教育主体身上并以与教育主体相协作的

方式参与人类对教育世界的创构，进而升级教育创新的路径。生成式人工智能是大脑器官的延伸，其组合创造、模板再生功能也只有通过放大教育创新功能来实现。有学者认为，生成式人工智能可以辅助学习者“利用更多的工具创造性地编辑、链接和整合知识，实现学习方式的变革”，这就是其参与教育创新的现实写照。

一是在认知前端发挥作用，即帮助教育主体对海量教育信息数据进行基于大模型的筛选与整合，以此来扩展教育主体的信息吞吐量，应对知识大爆炸时代的人类认知官能“无能”困局。无疑，借助 DeepSeek 等人工智能工具，可以从信息数据海洋中筛选出高相关、高价值、高密度的信息数据，并将之作为教育认识与教育创新活动的知识基座，进而克服了人类原生认知器官的功能阈限。有学者指出：“人工智能还是一个‘放大镜’，它只能根据信息量给予反馈：哪个观点的信息量多，它就放大哪个观点。”这种“放大”功能确保了教育创新始终与教育改革主流同频共振。

二是在认知中端发挥作用，即借助大模型、算法程序创新，以“创新合伙人”的身份介入教育创新进程。生成式人工智能为教育主体提供另一种认知思维方式，向教育创新实践植入另一种思维逻辑，推动教育主体从惯常性认识思维模式中走出来，以此加速教育创新思维的迭代升级。在生成式人工智能协助下，教育创新路径与逻辑不仅可以被轻松复制并将之推广到其他教育主体身上，还可以在程序员协助下创生出其他创新思维逻辑，进而从底层逻辑入手带动教育实践创变。

三是在认知终端发挥作用，即影响教育创新走向，系统提升教育创新的品质与成果。生成式人工智能参与教育实践之后，与自然人相并行的数字人、模拟人、孪生人诞生，每一个教育数字人、模拟人都可能成为教育主体历史经验的物理累积者，并将人的全部主体认知经验带入对新问题的思考与处理中。教育数字人、模拟人对教育创新的参与将缩短教育主体的迭代周期，进而大幅度提升教育创新的频率，缩小教育创新活动的代差。生成式人工智能技术不仅在塑造着教育创新主体，还将再造出教育数字人，延伸教育创造的功能。当前，人工智能技术正“以‘周’为迭代速度向前突进”，生成式人工智能将持续提升教育创新的平台与能级。

（二）催生双元复合式教育创新体

现代信息和传播技术的兴起和运用，对传统的以教师为中心、教材为中心、

教室为中心的教育教学模式产生了巨大的冲击，从根本上改变了人们对教、育、学的看法和做法。生成式人工智能延伸的不是人的眼睛、耳朵、嘴巴、手脚等普通器官，而是人的大脑、思维、心智，是具有一定自主反应能力的“准主体”，这就决定了必须将之作为一个与主体相对的“创新元”来看待。生成式人工智能的基本功能是知识生成性，未具备知识创生、观念创新的功能，其实质是基于“token”意义、计算逻辑的知识符号组合式再造。当模型规模大到一定程度时，生成式人工智能的应用效果就会显著提升，最终出现大模型中的能力涌现。随着生成式人工智能的加盟，教师的教育创新将步入“双元”组合、联手共创的新阶段，教育创新的疆域与空间被改写，一种以“人机协同、跨界融合与共创分享的新型发展模式”为特征的新型教育创新联合体出现。创新可分为原创型创新、组合型创新和模拟型创新三种。基于生成式人工智能的教育创新形态充其量只是一种组合型创新或模拟型创新，典型特征是“在已有4模型”的基础上增减、重组，从而获得新的功能”。它的入局会催生一种全新的教育创新架构，即“GAI+人”的复合式双元教育创新架构。更要进一步看，生成式人工智能的教育创新是基于智能体“数据捕捉分析反馈”的机械式线性创新，是依靠“大数据+大算力+强算法”实现的大规模数据组合创新，而教育创新则是基于真实问题情景与理性思维分析的自然创新或非线性创新，是借助头脑的顿悟力、想象力、批判力展开的小数据情景性创新；生成式人工智能教育创新在数据材料的初级组合创新中将展现出大数据、程序化、模型化处理的独特性能，发挥“1N”的创新优势，教育创新则在新观点涌现、新工艺研发、新方法锻造中发挥着“01”的创新优势。尽管两种教育创新类型、两个教育创新元在场景适用性、功能针对性、能动创造性等方面存在显著差异，但在具体问题解决中则具有较强的优势互补性、手段耦合性，生成式人工智能的教育创新功能可以作为人的“外脑”协助教育人开展教育创新实践，进而发挥“双脑合璧”的特异创新效能。生成式人工智能的加入使教育创新走上一条“机创”与“人创”双轨并进、双元互生的道路，大幅提升了教育创新活动的性能，突破了教育创新的视野局限与官能局限，增进教育创新活动的信息流，催生出一个划时代的双元复合式教育创新体，实现了教育创新形态的又一次突变。

（三）塑造新质态教育创新范式

当教育创新进入“双元协同创新”的时代，人的自然脑与人工智能的机器脑

同台竞技、携手联创，催生出了一系列全新的双元组合教育创新样式，标志着一种新质态教育创新范式的正式登场。生成式人工智能不仅具备了为师生定制教育教学工作方案、构建虚拟仿真学习场景、开发教育教学辅助图像视频、追踪教育教学活动状态、推送教育教学工作建议等多种功能，还具备了与师生主体开展专业对话、共创人机合作学习空间的新功能，一种逐步逼近“社会关系体验教育模式”的人机协同教育形态初见端倪。就目前而言，“立足教育大数据的人工智能，通过教育教学过程的数据采集、建模、智能分析和系统化的分析，实现教育教学决策的科学化、资源配置的精准化”，智能化思维与自然化思维携手攻关，将彻底改变教育创新的底层逻辑、运作机制与创新范式。

从教育创新逻辑上看，基于双重逻辑，即“程序思维 + 涌现思维”的教育创新逻辑出现，前者主要凭依的是程序开发思维、程序员思维与大模型思维，后者凭依的则是自然脑特有的情境性顿悟能力、格式塔生成能力、直觉顿悟力，二者联合行动，打破了单纯被人脑主控的教育创新逻辑。

从教育创新主体上看，基于双主体，即“数字人创新 + 教育人创新”的组合创新主体联合体诞生，在教育场景中隐现的数字孪生人，作为师生的另一种生命存在形态，将承载着师生的几乎全部数字教育空间经验进入教育创新舞台，尤其是基于大数据、大模型、大算力的高性能信息输入式创新必将充分彰显生成式人工智能赋能教育创新的功能优势，与之相应，真正的教育人则借助数字人的功能优势提升教育创新的知识基础与思维广度，教育创新将在“数字经验”与“实践经验”两个维度上跨界衔接、融创生成。

从教育创新网络上，一个巨大的教育创新网络正在形成，在元宇宙空间构建的教育数字人创新集群、师生在数字学习空间构建的自然人创新集群、教育行业生成式人工智能供应商构建的教育大模型创新集群、程序员与师生在合作开发应用中构建的跨界教育创新集群等，都将在教育创新任务上实现紧密协作、异质协同，构筑起一个巨型教育创客网络、教育知识生产集群，人类教育势必真正步入人机深度协作、共建共创共享的新发展阶段。

三、生成式人工智能时代教育创新范式走向追踪

生成式人工智能时代是人工智能充分赋能教育大发展的时代，是人的创新与

机器创新交融发展的教育时代。教育创新的新范式是：数字人与教育人两个教育创新元经由叠加、复合、协同、互联等途径实现携手共创，“人机”将共生出一系列教育创意，加快教育创新的节奏。

（一）走向双元叠加创新，带动教育创新功能升级

“双元叠加创新”是生成式人工智能时代教育创新的新形态、新样式。如何实现人脑与机脑间的深度叠加与复合创新，力争在混合创新中提升教育创新的能级与代次，是当今时代教育创新要考虑的问题。有学者指出，“迈向人机共教，是实现我国教育高质量发展的有效路径，是创造人类文明新形态的必然选择”，人机共教、人机互教、人机共育正是“双元叠加创新”的具体实现路径，沿着这条教育创新范式转型之道推进“GAI+ 教育创新”，势必带动教育创新实践的系统性提质。

一是推进“双脑”联动、“人机”融创，提升教育创新的功能。所谓“人机”融创，就是在藤别人脑与机脑各自创新机理、创新性能、创新专长、创新短板的基础上找到二者间的耦合点、互补点、共生点，据此建设人机共创教育平台、开发人机共学共教项目、拓展人机共生教育实践领域，进而实现机脑的组合创新与人脑的原始创新在具体教育实践课题上的联合联动联创。比如，当前可以借助豆包、元宝、深度探索等生成式人工智能工具，协助教师开发生动鲜活的动漫学材，开发人机共参、人机竞赛的学生作文教学课程，充分发挥智慧教育赋能教育创新的功能优势，让教育教学活动变得更有生气与活力。

二是实行人机教育创新分工，由机器人完成低阶教育创新、教育人完成高阶教育创新，在教育创新实践中推行专业化分工。现阶段生成式人工智能还处在弱人工智能发展阶段，其智慧水平仅限于低阶思维活动代理，其素材生成水平远未达到人类特有的“知识创生”高度，这就决定了教育类专用生成式人工智能的思维加工功能有限，交由其完成数据资料初级加工、静态教育数据处理加工、教育现象素材模拟式再制等低阶任务较为符合时宜。那些基于情景理解、批判想象、灵动思维的高创造性教育工作则必须交由人脑去亲自完成，以此实现人脑与机脑间的科学分工与分布协作。

三是优化教育创新流程，构筑“前端智能工具预创新 + 后端人脑真创新”的

教育创新新流程、新格局，梯次提升教育的综合创造力。生成式人工智能的优势在于其预训练模型的优势，预先借助生成式人工智能开展前端教育创新，在统观已有教育创造成果、进行分类处理加工、生成机器创新成果基础上再开展基于人脑的真正教育实践创新，生成新质教育成果、作品、知识，促使机器人教育创新与人脑教育创新接力式展开，构筑全新的人机教育创新“连续体”，一定是人工智能时代推进“双元叠加创新”的科学路径。

（二）走向多层级复合创新，缩短教育创新迭代周期

“GAI+”教育不仅有横向交叠式组合路径，还有纵向进阶式组合路径，这就是多层级复合。这一新的教育复合体将推动教育创新走上一条“多层级复合创新”的路子，进而大幅度缩短教育创新的生命周期。在这一梯级式多层级复合教育创新链路中，程序模型创新、数字人创新、教育人创新构成了教育创新阶梯的三个层级，彼此联动互促，将教育创新实践逐步推向新的高度。

其一是程序模型创新。这是“GAI+”教育创新的最底层建筑。这一教育创新层级的核心操作是：借助生成式人工智能的底层大模型研发，将已有的教育创新思维固化其中，诱发各种类、各水平教育创新思维的叠加效应，实现对现阶段教育思维智慧的综合吸纳与整体迭代。当前，生成式人工智能性能迭代提速，教育人的更多创新思维细节将被智能体模拟并实现，教育类人工智能与通用人工智能的深层创新逻辑都将被改良，其智慧水平正在一步步接近人类。在这种形势下，生成式人工智能将具备更强的教育创新辅助功能，教育创新的底层逻辑运行将更具活力与效能。

其二是数字人创新。这是“GAI+”教育创新的主体层装置。在人工智能时代，每一个人不仅有自己的肉体生命、精神生命，更有自己的数字生命，它以教育数字人、教育模拟人为物质载体再造了一个“虚拟主体”，师生的教育数字人完全可以借助教育人的智能学伴、智慧助理来经营自己的数字生命，绘制自己的数字画像，生成自己的数字人格。在人工智能世界中智能体“通过实时收集和智能分析学习数据，可以对学习过程及时干预、迭代、优化，并支持伴随式评价、综合素质评价和个性化培养”，与这一智慧教育过程相伴生的是师生教育数字人的持续演化，它将适时帮助师生耦合自己的原创性新经验与过往旧经验，使师生的教育生命与创新实践得以在数字世界中延绵生长。

其三是教育人创新。这是“GAI+”教育创新的灵魂层组件。智能时代下的教育必然从工具思维走向原点思维，将不可被人工智能替代的素养与能力作为培养的核心目标。在“GAI+”教育创新系统中，教育人不仅要发挥自己灵动性、涌现式、创生型教育创造的效能效力，还要善于用原始创新思维来规划、规训生成式人工智能体的知识、素材、作品的生成方向，整合其智慧生成功能，尤其是要立基上述两级知识创新实践，在融创式教育实践中提高教育创新的基点，充分发挥基于大场景、大模型、大数据的原始教育创新的新优势，实现教育知识能力的指数级增长。

（三）走向高仿真创新，大幅度提升教育创新的现实可行性

在生成式人工智能时代，人工智能不仅改变了教育创新的手段、工具、思维、逻辑，还提高了教育创新的逼真度与现实性，提升教育创新的品质与效度。在前人工智能时代，教育创新必须遵循的逻辑法则是合规律、合道德、合逻辑，而在后人工智能时代，教育创新的第四条法则是现实可行性或实践可转化性，这是因为人工智能完全可以将教育创新成果置于虚拟世界中进行仿真性模拟、仿真性推演、仿真性成效预判，过滤掉那些“天马行空”式教育创新和虚幻无根型教育创新，大幅度提升教育创新的成果转化率。

随着仿真实验的科学推进与逐步完善，人文社会科学研究成果的模糊性、主观性与柔软性等缺陷将部分被克服，研究结论的客观性随之逐步提升。生成式人工智能向教育创新实践的挺进，必然使仿真教育实验成为可能，研究者完全可以将大批量的数字人置于虚拟教育实验场景中，开展仿真教育实验、收集研究数据、预测研究结果、开展反馈调适，确保大部分教育创新成果具有现实可行性。

首先，借助生成式人工智能辅助研究者生成教育创意，提出可行性教育创新工作方案。无疑，生成式人工智能完全具备按照教育工作者的提示词去生成一系列教育改革创意的功能，教育人在参考借鉴这些教育创意的有效成分之后再启动教育创造、教育变革，进而可能生成更具科学性与可能性的教育创新工作方案。尽管人工智能体生成的教育创意不一定是最优化的，但它一定高于一般教育研究者、教育实践者的创新水平。只要研究者善于识别其中的关键创新点，并在“扬弃”精神指引下展开接续创新的探究活动，一种更趋完美的教育研究工作方案将会最终形成。

其次，开展基于生成式人工智能体的教育场景模拟，部分预知教育创新方案的预期效能。在本阶段，教育研究者完全可以签约教育主体被试的数字人、模拟人，将之置于教育教学数智实验场景中开展虚拟仿真实验，生成、收集、分析各项实验数据，对教育实验假设的可能后果进行全方位验证，从中发现问题、疏漏与缺陷，反向修正教育假设，确保教育创新成果的轮廓逐步清晰、数据逐步逼真、效果逐步现实。在仿真教育实验中，每个教育被试都可能允许自己的模拟人代理参加实验，被试模拟人在数智实验场景中的行为反应可减少主观性，更便于开展大规模实验，开展更为精准的大数据分析，从而克服被试亲身参与实验的诸多现实难题。最后，推进“仿真模拟+数据循证+多轮实验”研究设计，提升教育创新成果或方案的效能稳妥性与可行性。在未来教育创新中，“仿真模拟+数据循证+多轮实验”研究设计将成为教育研究创新、教育实践创新的标配，每一次教育创新实践都将“仿真实验、循证研究、多轮探索”纳入研究规划中去，都将得到虚拟数字场景验证数据的支持，大幅度提升教育创新实践的成功率、可信度与转化力。基于模拟人的教育仿真实验设计具有常规教育实验所不具备的优势，堪称对传统教育研究设计方案的一次颠覆性升级：它有效整合了循证研究、行动研究、实验研究的优势，可以借助虚拟人部分开展常规调查研究。在信息数据隐私保护的情况下将教育主体被试在模拟人身上保留的教育信息深度开发应用，将更多的教育研究变量考虑在教育研究与探索之中，无疑是进一步逼近教育生活真实状态的一种实验研究形态。

作者信息

龙宝新 陕西师范大学教师教育协同创新中心主任，教育学部教授、博士生导师

张 銓 陕西师范大学教育学部博士研究生

AI 赋能教育创新 传承行知思想

——苏州市教育学会、陶行知研究会专项培训圆满落幕

春和景明，研学正当时。为深入贯彻教育数字化转型战略，推动陶行知教育思想在智能时代的创造性转化与创新性发展，2026年3月22日至25日，由苏州市教育学会、陶行知研究会联合组织，华中师范大学陶行知国际研究中心承办的“AI赋能教育教学创新”专项培训在华中师范大学顺利举办。来自苏州市教育学会、陶行知研究会及各县市（区）的会长、秘书长与骨干教师等成员共36人齐聚江城武汉，开启了一场思想与技术交融的研学之旅。

开班启智，锚定教育数字化转型航向



3月23日上午，专项培训开班典礼在华中师范大学文科楼会议室隆重举行。

长江教育研究院副秘书长、“生活·实践”教育共同体秘书长、华中师范大学国家教育治理研究院院长助理、博士生导师刘来兵教授致欢迎词，秘书处相关人员以及全体参训学员出席典礼。苏州市教育学会常务副会长褚天生代表主办方致辞，为本次培训明确了目标方向、注入了精神动力。



褚天生常务副会长在讲话中指出，2026年作为“十五五”规划开局之年，国家正以前所未有的力度推进教育数字化战略，这标志着AI技术已从教育改革的“选修课”升级为“必修课”，教育工作者必须以时不我待的紧迫感主动拥抱变革。他强调，苏州市教育学会、陶行知研究会始终以“服务教育改革、引领教师发展”为核心使命，多年来深耕叶圣陶“教是为了不教”、陶行知“处处是创造之地，天天是创造之时，人人是创造之人”的教育思想与现代教育实践的深度融合。此次选择与华中师范大学——这所国家教育数字化战略重要智库携手，正是希望借助其雄厚的学术力量与实践经验，帮助苏州教育工作者找准AI赋能教育创新的精准支点，实现智能技术“算法”与育人初心“情怀”的深度融合。同时提升在数字化转型背景下服务基层、服务学校、服务教师的专业能力，让苏州本土教育家精神在智能时代焕发新的生命力。他预祝本次培训圆满成功，期待全体学员学有所获、思有所得，将学习成果转化为推动苏州教育高质量发展的强劲动能。

专家领航，解码AI与教育融合新路径

本次培训聚焦理论引领与实践考察双维度，特邀国内教育专家及数字化转型领域顶尖专家倾情授课。



3月23日上午，中国教育学会副会长、长江教育研究院院长、华中师范大学国家教育治理研究院院长、“生活·实践”教育共同体总负责人、陶行知国际研究中心主任周洪宇教授以《“三师课堂”：何谓、为何、何为》为题，深度阐释了“教师、小先生、AI智能师”三元协同课堂的新形态。周教授指出，“三师课堂”是人工智能时代中国本土原创的教学范式，其核心在于赓续陶行知“教学做合一”“小先生制”的教育智慧，通过明确AI与教师的角色分工，构建“人主机辅、技术赋能、情感联结”的良性教育生态，为破解当下教育规模化与个性化的矛盾提供了中国方案。





3月23日下午，参训学员赴武汉经济技术开发区神龙小学（湖畔校区）开展考察调研，与湖北省武汉经济技术开发区（汉南区）教育局党委委员、神龙小学教育集团总校长韩瑾围绕“三师协同·跨界融合－重构学习生态”开展深度座谈，韩校长结合学校鲜活案例，分享了AI技术在课堂教学、家校协同等场景的创新应用，让学员们直观感受了技术赋能基础教育的实效。



3月23日晚上，华中师范大学人工智能教育学部博士生导师上超望教授带来《人工智能赋能下的教育数字化转型与课堂教学创新应用》专题报告，从教学流程重塑、教师角色转型、智能工具实操等角度，系统解析了AI如何助力教学从标准化走向个性化，并通过生成式AI在课题研究、教学设计中的实操演示，为学员提供了可落地的技术应用路径。

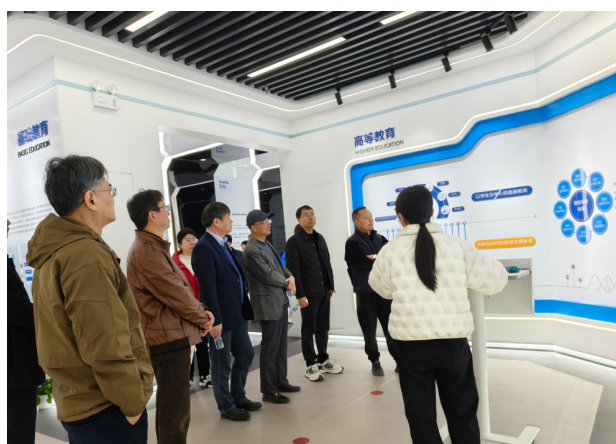


3月24日上午，刘来兵教授以《以教育家精神赋能教育数字化转型——基于数字化转型下的育人创新实践：师生共同课堂》为题，深入探讨了陶行知“生活即教育”思想与AI技术的逻辑契合点。他强调，智能时代的教育创新需坚守育人本质，避免技术异化，通过“师—机—生”三元主体协同，让教育回归真实生活、立足真实实践。

实地研学，感受科技与人文共生魅力



除高端学术讲座外，培训还精心安排了实地观摩环节，让学员在沉浸式体验中深化认知。3月24日，学员们走进华中师范大学校史馆与博物馆，在泛黄的史料与珍贵的文物中，探寻华中师大百年办学底蕴与教育思想传承脉络，感受“忠诚博雅、朴实刚毅”的校训精神与陶行知教育思想的精神共鸣。





3月25日上午，参训学员赴华中师范大学国家大数据实验室展厅开展实地观摩。作为国内领先的教育大数据研究平台，实验室工作人员详细介绍了教育大数据标准规范体系、智能教学平台研发及“小雅智能教学平台”应用等创新成果，展示了大数据技术在学情诊断、个性化推送、教育治理等领域的实践应用，让学员们直观领略了AI技术如何为教育精准化、个性化发展提供强力支撑。

知行合一，筑牢教育创新实践根基

四天的培训内容紧凑充实、形式丰富多样，既有前沿理论的深度解析，又有一线实践的经验分享，更有科技成果的直观体验。本次培训打破了传统教育思维的局限，搭建了交流学习的优质平台，不仅系统学习了AI赋能教育教学的前沿理念与实操方法，更深刻认识到陶行知教育思想在智能时代的强大生命力。

“周洪宇教授的‘三师课堂’理论让我们豁然开朗，原来AI不是要取代教师，而是成为教学的得力助手，这与陶行知先生‘教学做合一’的理念不谋而合。”大家一致表示，将把培训所学转化为实际行动，推动区域内教育数字化转型走深走实。

此次培训是苏州市教育学会、陶行知研究会落实教育强国建设要求的具体举措，也是深化陶行知教育思想当代实践、提升教育骨干数字素养的重要探索。通过理论学习与实地研学相结合的方式，为苏州教育系统注入了AI赋能教育教学创新的强劲动能，为推动苏州教育高质量发展、打造教育数字化转型标杆奠定了坚实基础。

周洪宇院长受邀出席《智慧教育》创刊暨资源中心期刊编委会成立座谈会

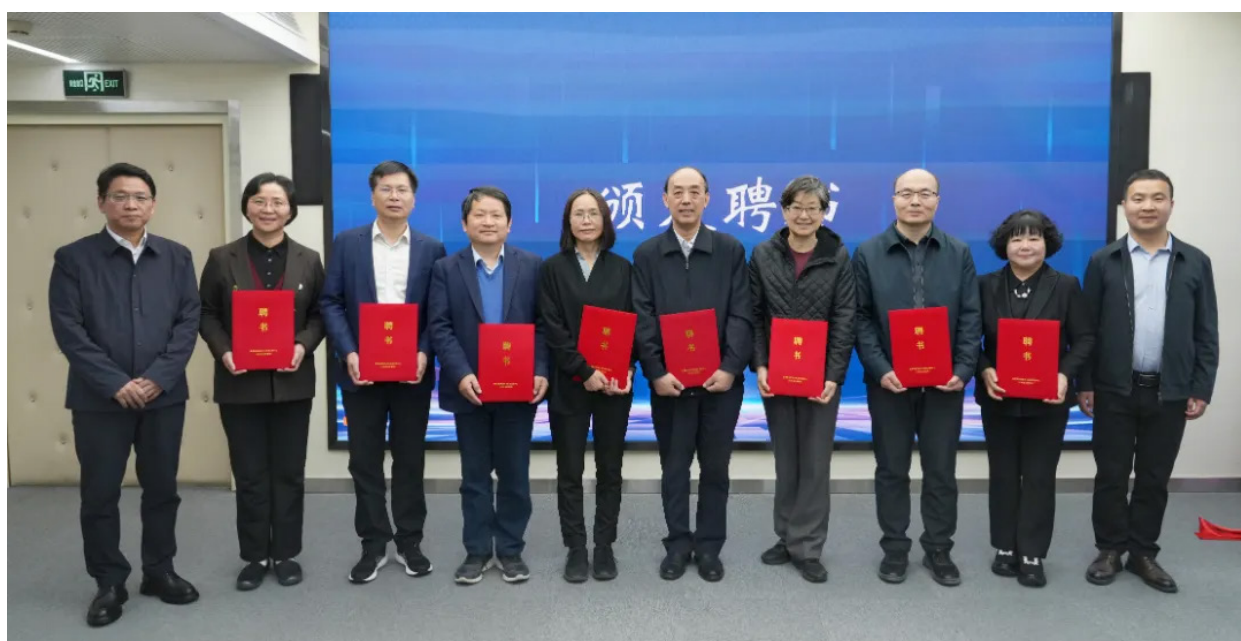
3月27日，教育部教育技术与资源发展中心（中央电化教育馆）在京召开《智慧教育》创刊暨教育部资源中心期刊编委会成立座谈会。教育部科技司信息化处任昌山处长、教育部教育技术与资源发展中心肖航副主任，来自北京大学、浙江大学、北京师范大学、华东师范大学、华中师范大学、东北师范大学、华南师范大学、首都师范大学等单位的编委会代表，内蒙古、浙江和成都七中初中学校等电教装备战线和学校代表，《人民日报》《光明日报》《中国教育报》等媒体代表，点猫科技等行业代表30余位专家学者出席会议，并围绕“数字教育研究暨资源中心期刊高质量发展”展开研讨。周洪宇院长出席会议并被聘为资源中心期刊编委会委员。



《智慧教育》期刊于2025年底获批，是国内首本关注国家智慧教育公共服务的创新性普及类刊物，将聚焦真实场景、回应一线关切、展示真实案例，致力于

打造教育工作者共同的“思想策源地”和“实践加油站”，努力成为一线教师可参考的实践指南、学校管理者可借鉴的决策参考、成为推进教育数字化的行动智库、展示教育技术装备创新成果的协作平台，同时向国际传播中国智慧教育公共服务的实践经验与独特贡献。期刊主管单位代表教育部科技司信息化处任昌山处长与期刊主办单位资源中心肖航副主任共同揭幕了创刊号。

会议当天举行了资源中心期刊编委会成立仪式。期刊编委会是资源中心加强期刊编辑出版工作，进一步发挥期刊学术平台作用的重要举措。首届编委会汇聚了国内智慧教育领域的十余名知名专家学者，他们将为期刊的选题策划、质量把关、方向指引提供强有力的学术支撑。会议上为编委代表颁发了聘书。



在“数字教育研究暨资源中心期刊高质量发展座谈会”上，《中国电化教育》《教育与装备研究》《智慧教育》三刊负责同志分别介绍了办刊情况。与会专家围绕期刊如何助力数字教育研究、服务国家教育数字化战略行动展开深入交流。专家们围绕提升期刊学术影响力、处理理论与实践关系、服务一线实践、加强国际传播等方面建言献策。

会议认为，当前人类社会正处于智能时代的井喷式爆发时期，智慧教育作为教育数字化转型的目标形态，不仅是技术和工具的革新，更是人类对美好教育的

共同求索。《智慧教育》的创刊恰逢其时，是立足发展新阶段、贯彻发展新理念、构建教育出版新格局的重要举措，为立体化宣传报道中国数字教育学术研究和实践创新成果提供了又一广阔的平台。期刊选择在国家智慧教育公共服务平台上线四周年之际创刊，体现了融入国家智慧教育公共服务体系建设责任担当，意义深远。

会议同期举办了资源中心期刊编委会第一次工作会，编委们认为，教育部资源中心承担推动现代技术的教育应用和数字资源的开发建设等多重职责，具有丰富的学术期刊办刊经验，主办的《中国电化教育》《教育与装备研究》学术期刊在数字教育领域具有重要影响力，《智慧教育》与其他两本期刊各有侧重、相互支撑，共同构成了资源中心服务教育数字化战略的期刊矩阵。编委会会积极发挥学术指导和引领作用，积极为期刊发展出谋划策。在未来发展中，希望三刊要资源共享、优势互补，深化新媒体运营与学术活动融合，让期刊成果更好服务一线、惠及更多教育工作者，为教育强国建设作出新的更大贡献。

周洪宇院长受邀出席第三届教育改革与发展论坛暨《河北师范大学学报（教育科学版）》新一届编委会会议

4月9日，由河北师范大学主办的“第三届教育改革与发展论坛暨《河北师范大学学报（教育科学版）》新一届编委会成立大会”在雄安新区顺利召开。中国教育学会副会长、长江教育研究院院长兼华中师范大学国家教育治理研究院院长周洪宇教授受邀出席会议。

会议由河北师范大学党委常委、副校长曾智安主持，河北师范大学党委副书记、校长李桂君出席大会并致辞。会上，主办方回顾了《河北师范大学学报（教育科学版）》的发展历程与办刊成效，详细介绍了学报在栏目建设、转载引用、学术影响及平台运营等方面的成果，同时汇报了河北师范大学教育学学科建设的成效、现存不足及未来展望，明确了学科与学报双向赋能、融合发展的总体思路。

周洪宇教授及其他与会专家充分肯定《河北师范大学学报（教育科学版）》创刊28年来取得的显著成绩，高度评价学报在教育史专栏建设、学术成果传播、青年学者培育等方面的突出贡献，并围绕坚守办刊初心、把握学术前沿、坚持守正创新、突出特色品牌、凝聚学术队伍、学报学科融合发展、服务国家战略等方面，提出了一系列具有针对性与操作性的建议。

会上，周洪宇教授受邀担任《河北师范大学学报（教育科学版）》新一届编委会委员。研究院也将以此为契机，持续加强与国内高校、学术期刊的交流合作，推动教育学科融合发展。

周洪宇院长受邀参加教育强国建设重大问题研究与中国教育学体系建构学术研讨会

4月11日，教育强国建设重大问题研究与中国教育学体系建构学术研讨会在北京师范大学召开。教育部党组成员、副部长、总督学王嘉毅，中国高等教育学会会长林蕙青，北京大学党委书记何光彩，同济大学党委书记、中国工程院院士郑庆华，北京师范大学党委书记程建平，北京市教育委员会主任李奕，北京师范大学资深教授顾明远，吉林大学资深教授张文显，华东师范大学教育学部主任、终身教授袁振国等出席开幕式。开幕式由北京师范大学党委常委、副校长康震主持。长江教育研究院院长兼华中师范大学国家教育治理研究院院长周洪宇教授应邀出席会议，并参与第二场圆桌会议研讨。

王嘉毅副部长在讲话中指出，我国正处于“十五五”规划奠基开局和教育强国建设全面攻坚的关键阶段，教育学自主知识体系作为教育部统筹布局的29个一级学科重大专项的重要组成部分，既是一门学术体系，也是一套育人体系，更是一个支撑教育决策的政策咨询体系。

北京师范大学党委书记程建平强调，今年是习近平总书记在哲学社会科学工作座谈会上发表重要讲话十周年，本次研讨会既是对“加快构建中国特色哲学社会科学学科体系、学术体系、话语体系”重要指示的深入贯彻，也是对教育强国战略的积极响应。

中国高等教育学会会长林蕙青在致辞中强调，加快构建自主知识体系，是回应中国教育重大改革实践问题的迫切需要，是提升教育学科主体性和理论自信的迫切需要，是增强中国教育国际传播能力和国际学术影响力的迫切需要。

在大会报告环节，何光彩、李奕、顾明远、张文显、郑庆华、袁振国等六位专家分别作主旨报告，围绕教育学自主知识体系建构、教育综合改革、中国共产党教育价值观、人工智能与未来教育等议题展开深入论述。

在第二场圆桌会议环节，周洪宇教授与北京师范大学教育学部部长朱旭东、北京外国语大学原党委书记王定华、清华大学教育学院院长石中英、北京大学教育学院院长蒋凯等学者，围绕“中国教育学自主知识体系建构：理念、成果与高质量传播”主题展开深入对谈，分享学术观点。周洪宇教授就“在贯通古今、融汇中西中构建中国教育学自主知识体系”为主题展开交流，他表示，构建中国教育学自主知识体系是一项关乎方向、关乎根基、关乎长远的系统工程，需要学界同仁协同攻关、开放共享，共同推动中国教育学学科体系、学术体系、话语体系建设。



周洪宇院长受邀出席红色教育数智记忆 实验室建设与发展研讨会

4月15日，红色教育数智记忆实验室正式建立暨江苏省基础教育前瞻性教学改革实验项目、江苏省基础教育教学成果重大储备项目“入脑入心：以文化记忆推动红色资源育人的区域实践”研究推进会在南京晓庄学院隆重举行。第十三届全国人大常委会委员、中国教育学会副会长、长江教育研究院院长兼华中师范大学国家教育治理研究院院长周洪宇教授，南京晓庄学院党委书记张策华、党委副书记易永建，省教育厅社政处副处长陈磊，市教育局党组副书记、市委教育工委专职副书记柴耘，市教育局宣德处处长王京，浙江外国语学院副校长柴改英，南京大学信息管理学院教授、“雨花台烈士纪念馆——南京大学国家革命文物协同研究中心”主任李刚，南京师范大学马克思主义学院教授王跃等出席研讨会。全市德育干部、德育科研员、记忆德育实践学校、记忆育德专项课题主持人代表近300人参加活动。

红色教育数智记忆实验室成立仪式

活动的第一阶段是红色教育数智记忆实验室成立仪式，由南京晓庄学院党委副书记易永建主持。



南京晓庄学院党委书记张策华致辞，他提出作为陶行知先生亲手创办的学府，拥有深厚的红色底蕴，运用数智技术激活红色资源、创新育人模式，既是时代要求，更是义不容辞的责任。红色教育数智记忆实验室的成立，是主动回应教育数字化转型战略需求、系统推进红色资源数字化、红色教育智能化的重要举措，对推动红色文化传承与现代信息技术深度融合、落实立德树人根本任务具有重要意义。实验室在未来的建设与发展中，要突出红色内核、强化数智赋能、注重开放共享、坚持成果导向、健全保障机制，努力建成传承红色文化、引领教育数字化转型的亮丽名片。



在成立仪式上，柴耘和陈磊为“红色教育数智记忆实验室”揭牌，周洪宇和张策华为实验室执行主任刘大伟、贾学军，副主任李志兵、程功群颁发聘书，大家共同见证了实验室正式成立。





周洪宇院长在讲话中强调，实验室建设要立足红色基因传承与数智技术创新，聚焦红色资源数字化开发、沉浸式红色教育场景构建、记忆育德理论研究与实践推广等重点工作，努力打造集教学、研究、实践于一体的高水平育人平台，让红色基因代代相传，让行知精神焕发新的时代光彩。

红色教育数智记忆实验室建设与发展研讨

第一会场由南京晓庄学院马克思主义学院贾学军院长主持。研讨会上，南京教育科学研究所所长、研究员刘大伟做了“红色教育数智实验室建设举措”专题汇报。全体与会专家学者、省市教育主管部门领导，共同围绕实验室功能定位、资源建设、技术路径、产学研协同等核心议题展开深入研讨。大家认为，实验室立足南京、面向全国，站位高、起步实，前期基础扎实，建设方案切实可行；建议实验室要进一步强化内容整合、技术标准和应用场景三大支撑，加快构建全国性的红色教育数字记忆体系，力争早日建成省级乃至国家级文科实验室。





第二会场由南京市教育科学研究所组织全市记忆育德实践学校代表围绕学科记忆育德实践阶段成果凝练、未来如何做好实验室数据采集进行研讨。南京市栖霞区教师发展中心中学物理教研员、南京大学博士陈晨，建邺区教师发展中心小学英语教研员、正高级教师张丽宁，南京市齐武路小学校长、副书记邵丹丹，栖霞区教师发展中心初中语文教研员潘丹婧四位专家分别进行了“记忆育德：物理学科中的育人密码”“立德树人视域下英语记忆育德的实践探索专题分享”“跨学科融合，让记忆成为育德新载体”“文化记忆·语文在场：记忆育德的栖霞初中语文实践”进行专题分享。他们通过教材整合、学科教研建模、跨学科实践、文化思考，共同探析如何从价值、过程、关系等角度追问学科育人价值导向、记忆育德过程中知情意行统一、德育与教学的双向深度融合等学科育人突破锚点。





红色教育数智记忆实验室由我所刘大伟研究员提出构想并发起成立，这一成果既是南京市教育科学研究所持续进行江苏省前瞻性教学改革实验项目和江苏省基础教育教学成果重大储备项目研究的阶段性成果，也是主动与高校进行多方联动协同，落实《教育强国建设规划纲要（2024-2035 年）》中建好哲学社会科学实验室的相关精神与要求，努力从区域层面呈现数智时代德育理论与实践研究创新，构建可复制、可推广的红色资源育人“南京样本”。

周洪宇院长受邀出席“求是教育论坛”

4月18日,浙江大学教育学院第119期“求是教育论坛”在222会议室顺利举行。本期论坛以“陶行知研究与教育学自主知识体系建设”为主题,特邀第十三届全国人大常委会委员、中国教育学会副会长、中国陶行知研究会学术委员会主任、华中师范大学教育学院二级教授、国家教育治理研究院院长周洪宇教授担任主讲嘉宾。讲座由教育学院教育学系副主任李木洲长聘副教授主持,教育学院院长阚阅教授及部分硕博研究生参加了本次论坛。

讲座伊始,周洪宇教授结合自身学术背景,回顾了45载的陶行知研究心路历程。在章开沅等先辈学人影响下,他以严谨的历史视野与宏大的文化关照为基石,完成了从“被动结缘”到“主动深耕”的思想转变。他认为,陶行知研究不仅是个案研究,更是具备高度“文化主体性”的探索,其核心在于将中华优秀传统文化与现代教育理论深度融合。周教授将陶行知研究作为其治学治教的“学术根据地”,先后撰写并出版27部学术专著,构建起系统而坚实的理论体系。



随后，周教授分享了其从个案研究，逐步迈向更为宏阔的中国教育史与教育活动史研究的学术转向。他指出，传统教育史研究长期局限于思想史与制度史，导致教育过程中的“活动”维度被长期忽视。教育活动史的提出，正是为了打破史料获取的客观局限，应对当代教育改革回溯历史经验的现实需求。通过对生活史、身体史、情感史等微观领域的向内挖掘，教育史研究得以构建起一个宏观与微观相结合、理论与实践双向驱动的研究范式。此外，他积极倡导以全球史观和全球视野拓展研究边界，强调以深度了解国际学术界为前提，中国教育研究才能在跨文化对话中实现真正的回应、参与和引领。

在谈及学术研究的落地转化时，周洪宇教授介绍了新时代“马工程教材”通过构建三个“三位一体”的立体化结构实现的原创性突破。该结构打破了传统叙事的偏颇，打通了历史、现实与未来的时空逻辑，为构建中国特色、世界一流的教育学自主知识体系奠定了坚实的史学根基。结合自身经历，他进一步阐释了“参与史学”从理论走向现实的路径，强调教育研究不应止于纸笔，而应深度参与国家治理、政策制定与法律完善。从教育治理学丛书的出版到教育法典的编纂，研究者应当在治理现代化进程中发出中国声音，贡献中国智慧。



最后，周洪宇教授分享了教师－小先生－AI智能师“三师课堂”的创新实践。这一模式实现了陶行知“小先生制”在人工智能时代的创造性转化，通过教师、学生与AI的协同互补，探索未来教育的新形态。周教授强调，从“拿来化”到“特色化”“自主化”，中国教育学已进入第六代学人的探索期。教育研究应当坚持以社会需求为导向，在传承与创新中完成中国教育学自主知识体系的当代构建。

周洪宇教授的讲座立意高远、逻辑缜密，既展现了深厚的人文底蕴，又前瞻性地展望了智能时代的教育变革，引发了在场师生的强烈共鸣。在交流环节，周教授就“三师课堂”内涵及研究等问题做了进一步阐述。

讲座尾声，李木洲长聘副教授对讲座进行了精彩点评与总结。他表示，周洪宇教授不仅为我们系统梳理了陶行知研究的学术谱系，更指明了建构教育学自主知识体系的科学路径，这对提升我院科研水平与人才培养质量具有重要的指导意义。最后，他对周洪宇教授为我院师生带来的精彩讲座表示衷心感谢。



《教育治理研究》的征稿通知

随着国家对提高治理体系和治理能力现代化的高度重视，提升教育治理体系和治理能力现代化水平已日益成为教育改革发展的当务之急。为顺应此教育变革大趋势，促进高质量教育体系建设，创刊于 2015 年的《长江教育论丛》已从 2022 年起正式更名为《教育治理研究》，并顺利出版和上线知网。

《教育治理研究》（半年刊）由华中师范大学国家教育治理研究院、长江教育研究院联合主办，致力于成为以教育治理研究为主题和特色的集刊，入选中国人文社会科学（AMI）核心集刊。著名教育家、原中国教育学会会长、北京师范大学教育学部顾明远教授担任本刊顾问。朱永新、徐辉、张力、王定华等一批全国知名教育专家组成本刊编委会。

（一）主编介绍



《教育治理研究》主编周洪宇

周洪宇，第十三届全国人大常委会委员，中国教育学会副会长、中国教育发展战略学会副会长，长江教育研究院院长、华中师范大学国家教育治理研究院院长、教授。

（二）投稿须知

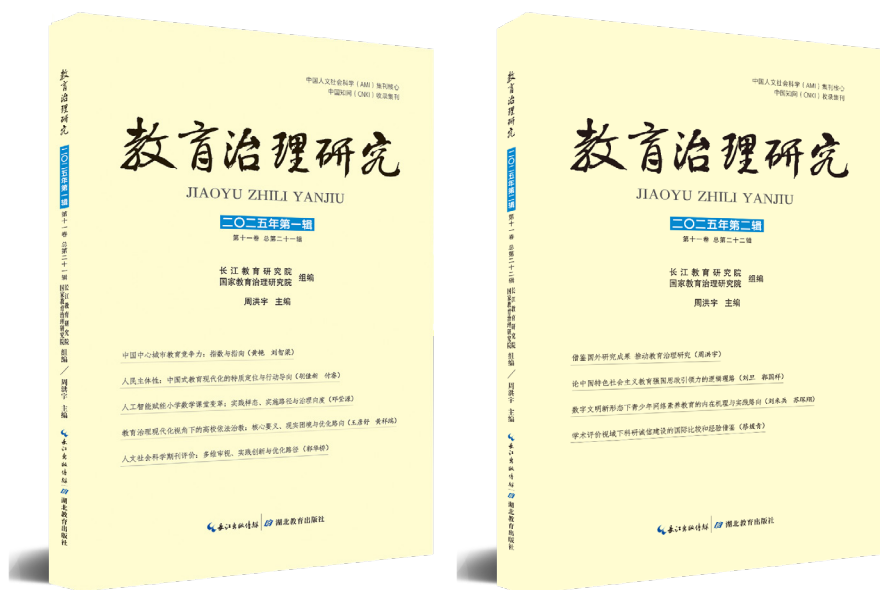
1. 征稿对象：关于基础教育、高等教育、职业教育等领域的知名学者、教师、专业研究人员、管理人员等。
2. 栏目设置：卷首语、特稿、习近平总书记教育思想论述、教育治理学研究、人工智能与教育治理、教育舆情监测与教育治理、基础教育治理、高等教育治理和有关专题等，出版时根据稿件内容进行适当调整。同时，我们特别设立“青年学者专栏”，致力于发表年龄在35岁左右的青年学者及优秀博士研究生的学术论文，每期力争2—3篇，以凸显本刊特色。
3. 内容要求：选题新颖、数据可靠。鼓励就教育治理领域某一主题进行深入、细致的研究，篇幅在1万字左右为宜。
4. 稿件来源：全国征稿，匿名三审。投稿邮箱为 jyzlyj@126.com。
5. 截稿日期：2026年第2期截稿日期2026年9月30日；2027年第1期截稿日期2027年3月30日。
6. 学术规范：请自觉遵守学术规范，勿一稿多投。本刊审稿周期为二个月，二个月未收到回复可另投其他刊物。

（三）体例格式

1. 基本要件：论文标题、作者署名及简介、摘要与关键词、正文、注释、参考文献。如涉及资助项目，请注明项目来源、名称和编号。
2. 论文标题（不超过20字，中英文）。
3. 作者署名及简介：作者署名一般不超过4人；作者简介包括姓名、性别、籍贯、所获学位、任职机构（正式全称）、职务/职称。
4. 摘要与关键词（中英文）：摘要须独立成篇，完整、准确地概括文章的实质性内容；关键词一般不超过5个。且中英文应相互对应。

5. 正文：①标题一般分为三级，第一级标题用“一”“二”“三”等标示，第二级标题用“（一）”“（二）”“（三）”等标示，第三级标题用“1.”“2.”“3.”等标示；②图、表和公式均用阿拉伯数字连续编号，如图1、图2和表1、表2，以及（1）（2）等。图和表应有简短确切的题名，图号图名应置于图下，表名表号置于表上，公式号置于右侧。

6. 注释、参考文献：著录格式为“顺序编码制”，采用页下注形式，除电子文献外其余均需标记页码，作者需保证文献的真实出处，如有需要编辑部会请作者提供引文原文。具体格式，参照引用性注释采用《文后参考文献著录规则》（GB/T7714—2015）。



长江教育研究院是在湖北省教育厅的支持下，由华中师范大学和湖北长江出版传媒集团有限公司联合发起，于 2006 年 12 月 16 日成立的教育研究机构。由第十三届全国人大常委，四届全国人大代表（2003—2023），湖北省人大常委会原副主任、中国教育学会副会长、华中师范大学教授周洪宇担任院长。

长江教育研究院本着“全球视野、中国立场、专业能力、实践导向”的指导思想，“民间立场、建设态度、专业视野”的立院原则，聚集了一批国内外优质教育专家资源，搭建了一个以文化出版企业为依托、联系相关教育专家和教育管理部门的平台，形成了以学术研究为基础、政策研究为重点、出版企业为依托、政府支持和社会参与为支撑，“学、研、产、政、社”优势互补、协同推进的新型体制机制。

多年来，长江教育研究院一直致力于打造新型教育智库“重器”，努力让智库的“谋划”转化为党和政府的决策，智库的“方案”转化为实际行动，智库的“言论”转化为社会共识，更好地为改革奉献力量。自 2016 年来，连续三年在中国智库索引评选中社会智库类排名稳居前三。2017 年入选中国社会科学评价研究院“2017 年度中国核心智库”。



联系电话：027—87671389

官方邮箱：changyanyuan@qq.com

官网地址：<http://cjy.com.cn/>

公司地址：武汉市洪山区雄楚大道 268 号省出版文化城 B 座 6 楼