



学术探索
Academic Exploration
ISSN 1006-723X, CN 53-1148/C

《学术探索》网络首发论文

题目：硅基与碳基的共舞：知识生产革命中的学术创新逻辑——以传播学发展为分析对象
作者：喻国明
网络首发日期：2026-08-27
引用格式：喻国明. 硅基与碳基的共舞：知识生产革命中的学术创新逻辑——以传播学发展为分析对象[J/OL]. 学术探索.
<https://link.cnki.net/urlid/53.1148.C.20260827.1120.004>



网络首发：在编辑部工作流程中，稿件从录用到出版要经历录用定稿、排版定稿、整期汇编定稿等阶段。录用定稿指内容已经确定，且通过同行评议、主编终审同意刊用的稿件。排版定稿指录用定稿按照期刊特定版式（包括网络呈现版式）排版后的稿件，可暂不确定出版年、卷、期和页码。整期汇编定稿指出版年、卷、期、页码均已确定的印刷或数字出版的整期汇编稿件。录用定稿网络首发稿件内容必须符合《出版管理条例》和《期刊出版管理规定》的有关规定；学术研究成果具有创新性、科学性和先进性，符合编辑部对刊文的录用要求，不存在学术不端行为及其他侵权行为；稿件内容应基本符合国家有关书刊编辑、出版的技术标准，正确使用和统一规范语言文字、符号、数字、外文字母、法定计量单位及地图标注等。为确保录用定稿网络首发的严肃性，录用定稿一经发布，不得修改论文题目、作者、机构名称和学术内容，只可基于编辑规范进行少量文字的修改。

出版确认：纸质期刊编辑部通过与《中国学术期刊（光盘版）》电子杂志社有限公司签约，在《中国学术期刊（网络版）》出版传播平台上创办与纸质期刊内容一致的网络版，以单篇或整期出版形式，在印刷出版之前刊发论文的录用定稿、排版定稿、整期汇编定稿。因为《中国学术期刊（网络版）》是国家新闻出版广电总局批准的网络连续型出版物（ISSN 2096-4188，CN 11-6037/Z），所以签约期刊的网络版上网络首发论文视为正式出版。

硅基与碳基的共舞：知识生产革命中的学术创新逻辑

——以传播学发展为分析对象

喻国明

(北京师范大学 新闻传播学院, 北京 100875)

摘要：人工智能的深度介入正在重塑知识生产的底层逻辑，一场从“工具革命”到“认知革命”的历史性转型已然展开。本文以传播学科的创新发展为经验场域，在喻国明教授“认知合伙人”研究纲领的基础上，提出“学术创新的终极目标在于推动知识边界的突破，而非维护人类认知的垄断性尊严”这一核心命题，并构建“双主体协同—目标位移—责任锚定”的三维分析框架。研究认为：其一，硅基智能与碳基生命构成能力图谱的深层互补，AI 在认知层面可以扮演主体角色而不必在意识层面取代人类，“认知不等于意识，智能不等于主体性”是双主体协同的本体论前提；其二，学术创新的价值坐标应从“人类中心主义”位移至“知识本位主义”，以知识边界的突破为第一性目标，人类学者从“经验执行者”升维为“逻辑架构师”与“意义裁决者”；其三，能力的让渡绝不意味着责任的让渡，价值判断不可算法化、知识合法性依赖社会建构、法律问责需要可追溯主体，构成责任不可让渡的三重硬约束。本文进一步论证：传播学科因研究对象与技术的高度嵌合，既是知识生产革命的风暴中心，也是人机协同制度探索的先行实验场。研究旨在为智能时代负责任的人机共治提供传播学视角的理论坐标。

关键词：人机协同；知识生产；双主体；学术创新；认知合伙人；责任伦理；传播学科

引言：当“省思时刻”到来

学术研究的工具谱系从未停止更新，但工具从未像今天这样逼近“主体”的边界。从望远镜延伸人类的视力，到计算机延伸人类的算力，工具始终作为“客体”服务于“主体”的认知需求。然而，当大语言模型 (LLM) 能够撰写理论综述、生成研究假设、甚至提出看似原创的学术命题时，一个根本性的追问浮出水面：当工具开始“思考”，谁才是知识生产的主体？必须指出，AI 的本质不是工具升级，而是人类社会“操作系统”的底层重构——“不是换个新软件，而是重构人机交互方式、资源分配规则与社会运行架构”。笔者曾将 AI 的演进划分为三个阶段：从“能看会认”的认知智能，到“能说会画”的生成式智能，再到“能做会动”的物理智能。^①当前我们正处于生成式智能向物理智能过渡的

基金项目：国家社科重大项目 (25&ZD177)

作者简介：喻国明，北京师范大学新闻传播学院教授、北京师范大学传播创新与未来媒体实验平台主任、中国新闻史学会传媒经济与管理专委会创会理事长。

^① 喻国明：《范式跃迁：AI 时代获取未来入场券的密钥》，《广西大学学报（哲学社会科学版）》2026 年第 1 期。

关口, AI 第一次在“认知操作”的层面与人类平起平坐——它能够完成以往只有人类才能完成的综述、推理、写作与创造。当“操作系统”被重装, 知识生产的底层逻辑便不再可能原样运转。换言之, 当 AI 开始思考, “谁来思考”便成为时代之问——而答案不在人机之间, 而在人机之中。

传播学科——这门以人类信息交往与意义建构为核心关切的知识领域——恰恰置身于这场革命的风暴中心。它不仅是被 AI 深刻改变的研究对象, 更是观察知识生产方式转型的绝佳样本。从计算传播学到智能广告投放, 从算法推荐机制到 AI 生成内容 (AIGC) 的传播效应, 传播学的每一个分支都在经历着研究范式的地壳运动。更重要的是, 传播学具有“研究对象”与“经验场域”的双重身份: 它既研究 AI 作为传播主体的全新现象, 又以自身学科的知识生产实践为样本, 检验人机协同的普遍逻辑。换言之, 智能时代的传播学科需要在“人—机”协同的框架下重新审视人与技术的关系^②——这一判断同样适用于整个知识生产体系。

面对这一变局, 学术界存在两种值得警惕的极端倾向。一种是保守的人类中心主义, 试图通过制度壁垒将 AI 拒之门外, 以维护人类认知的垄断性尊严; 另一种是激进的技术决定论, 认为 AI 将全面接管知识生产, 人类将退化为数据的附庸。本文认为, 这两种倾向均未能准确把握硅基智能与碳基生命关系的本质: 前者低估了 AI 作为认知伙伴的变革潜力, 后者高估了算法对意义世界的统摄能力。应该指出, 将传播者与传播技术彼此对立的“二分法”思维是根本错误的, “以人为本”标尺下的“人—机”协同才是人工智能时代的不二之选。这一论点为本文提供了基本立场: 人机关系不是“取代”与“被取代”的零和博弈, 而是“协同”与“共生”的正和演进。

本文的核心主张明确而坚定: 学术创新的终极使命是推动知识边界的突破, 而非守护人类认知的垄断性尊严。但这并不意味着对人类主体地位的消解, 而是指向一种更为深刻的分工——在能力的赛道上让 AI 尽其所能, 在意义的判准上让人守其必守。这一主张的提出, 既是知识生产革命的逻辑必然, 也是笔者在过去的论文中所体积的“认知合伙人”命题在学术创新领域的延伸: 当 AI 进化为能够进行复杂推理、创意生成甚至情感模拟的“认知合伙人”, 人类学者的核心能力便从提供“标准答案”转变为定义“关键问题”、设定 AI 探索的“价值边界”^③。

据此, 本文构建“双主体协同—目标位移—责任锚定”的三维分析框架: 第一, 勘定硅基与碳基的能力图谱, 论证“双主体协同”何以可能; 第二, 重估学术创新的价值坐标, 论证“目标位移”为何必要; 第三, 锚定责任的价值维度, 论证“责任不可让渡”如何可能; 第四, 以传播学科及其交叉领域的典型案例对三维框架进行经验检验; 第五, 将理论视野从个体知识生产提升至传播学科发展的宏观镜鉴。全文以传播学为分析对象, 旨在为智能时代的知识生产革命提供一幅既开放又审慎的理论图景。

一、能力图谱的深层互补: 何以走向“双主体”协同

在探讨人机关系时, 首先必须直面认知能力的比较。AI 的崛起并非简单的效率提升, 而是引入一种异质性的认知逻辑。这种逻辑与人类认知既存在冲突, 更存在深层的互补性。我们需要在认识论层面厘清二者的特质, 从而为“双主体”协同奠定本体论基础。

(一) 硅基智能的认知优势与结构性盲区

笔者的核心论点: AI 的“理解”是统计学的拟合, 而非生命体验的领悟; 它的“创新”是组合空间

^② 喻国明:《“以人为本”标尺下的“人—机”协同》,《新闻与写作》2022 年第 10 期。

^③ 喻国明:《在认知、连接与意义的交汇处重塑世界——数字化转型时代基于传媒业洞察的四大范式变革》,《青年记者》2026 年第 3 期。

的搜索,而非存在困境的突围。理解这一论断,是评估硅基智能学术潜能的前提。

从优势维度来看,硅基智能展现出了超越人类生理极限的认知特质。其一是海量记忆的并行调用能力。AI能够以毫秒级的速度扫描数以亿计的文献与数据,这种“全知”视角是人类学者受限于生理寿命与阅读速度所无法企及的。其二是高维数据的模式识别能力。在复杂变量的纠缠中,AI能够识别出人类直觉难以捕捉的隐性结构——这正是“AI for Science”浪潮的认知基础:从AlphaFold对蛋白质折叠的破解,到生成式模型对材料结构的预测,硅基智能正在以前所未有的方式拓展人类认知的疆域。其三是跨域迁移的联结能力。大语言模型在海量语料上习得的“世界知识”,使其能够在学科边界之间自由穿梭,为交叉学科的创新提供“知识桥接”。从本质上讲,AI对世界的理解是对世界逻辑的各种层次的价值连接、价值匹配——这种连接能力,恰是学术创新最稀缺的资源。

以传播学科为例,在数十年的学科发展中,积累了海量的碎片化成果。人类学者往往受限于阅读量与记忆偏好,难以把握学科全景。而AI可以通过对数万篇传播学论文的语义网络分析,自动识别理论演进的隐性脉络。例如,它可能发现“议程设置”与“框架理论”之间未被充分理论化的接缝,或者从长达三十年的受众研究数据中提炼出“媒介使用—社会信任”的非线性关联假设。这些发现未必直接构成理论创新,却为人类学者提供了高密度的“创新原料”。

然而,硅基智能的优势也恰恰注定了它的结构性盲区。首先是缺乏“肉身在场”所带来的具身认知。正如第二代认知科学所揭示的,身体不仅是心智的物理载体,更是认知的锚点。这一观点深植于梅洛-庞蒂的知觉现象学传统,他强调身体—主体是我们“在世存在”的基底,所有的抽象意义与高级认知均源于身体与世界互动的具身体验^④。当代具身认知研究进一步证实,人类的概念系统并非脱离肉体的符号运算,而是深深锚定于身体的感觉运动系统及其与环境的动态交互之中^⑤。传播学研究的核心——身体共在、情绪传染与场景意义,都依赖于具体的、肉身化的知觉经验。AI可以分析“在场”的数据,却无法理解“在场”本身。笔者将其概括为人类“三张底牌”中的第一张:具身体验的不可替代性——“AI能模拟‘看到’‘听到’,但无法复刻人类的‘触觉’‘嗅觉’,更无法替代‘身处现场’的情感共鸣”^⑥。这并非技术的阶段性缺陷,而是认知的构成性差异。

其次,价值中立表象下隐藏着算法偏见。AI的“客观”计算建立在对既有训练数据的拟合之上,而既有的数据本身就是社会结构性不公的产物。例如,在基于社交媒体数据训练的情感分析模型中,少数族裔或边缘群体的语言习惯往往被误读为负面情绪。如果不加批判地采纳AI的分析结论,学术研究便可能在“客观”的外衣下再生产社会偏见——这正是后文案例三所揭示的“算法偏见”问题的认识论根源。

最后,大语言模型存在的“幻觉”问题,在认识论层面揭示了其局限。逻辑自洽不等于事实真实,统计合理不等于意义有效。AI可以在逻辑链条上完美推导出一个错误的结论,这在学术创新中是致命的。我们在论及“批判性算法审计”能力时应特别强调,AI的输出并非“绝对可靠”,可能存在数据偏见与逻辑漏洞,研究者必须具备“审查能力”。一句话,硅基智能拥有无限的计算与检索能力,却永远缺少一枚“肉身”的锚——而这枚锚,恰是人类认知的起点。

(二) 碳基生命的认知瓶颈与不可替代性

相较于AI的“无限”,人类碳基生命面临着显而易见的生理性瓶颈:工作记忆容量有限(米勒的 7 ± 2 理论)、认知负荷阈值、以及注意力的稀缺性。这些生物学上的限制,曾促使人类发明各种工具

④ 梅洛-庞蒂:《知觉现象学》,姜志辉译,北京:商务印书馆,2001年,第10-15,160-168页。

⑤ 叶浩生:《具身认知:认知心理学的新取向》,《心理科学进展》2010年第5期。

⑥ 同①

来延伸自我,但也规定了人类认知的边界。然而,正是这些局限,恰恰构成了人类认知不可替代性的另一面——意义感的来源。

这种意义感体现为一种将碎片信息整合为价值整体的直觉性判断力。首先是提问的原创性。提出“值得研究的问题”本身是一种价值判断,而非单纯的信息处理。AI 可以基于提示词生成一百个符合学术规范的研究问题,但只有人类学者能够基于对社会现实的敏锐洞察,判断哪一个是“值得的”、哪一个是“切肤的”。比如,传播者的核心能力将从提供“标准答案”转变为定义“关键问题”、设定 AI 探索的“价值边界”——提问的原创性,正是“定义关键问题”的能力内核。一个精准的问题,抵得上一百个泛泛的答案;而“问什么”的权利,始终是人类学者的第一权利。

其次是直觉的跃迁性。库恩在《科学革命的结构》中指出,科学范式的转换往往源于某种不可言说的科学直觉。这种直觉类似于波兰尼所说的“隐性知识”,它不是规则化的程序,而是一种基于长期积累的顿悟^⑦。在传播学研究中,当一位学者从复杂的受众行为数据中突然领悟到“媒介即按摩”式的洞见时,这种跃迁无法被还原为算法步骤——它是“知其然亦知其所以然”之外的“知其所以不可言说”。

再次是意义的建构性。传播研究的核心关切——意义如何在符号交换中生成、维系与转化——本质上是对“人的存在方式”的追问。这需要研究者以“人的体验”为前提。例如,对深度伪造传播效果的研究,AI 可以精准测量受众认知变化的数据曲线,但只有人类学者能深切理解“当一个人无法确认所见是否为真”时的存在性焦虑,并将其转化为有温度的理论范畴。所谓“意义生产的稀缺性”——AI 能处理信息,但不能生产意义——在此获得最充分的体现:在信息爆炸的时代,人类是“炼金术士”,能从海量碎片中提炼出有价值的观点、有温度的叙事,这是算法的“效率逻辑”永远无法抵达的维度。显然,人类的能力边界,恰恰是人类意义的起点——正因为“有限”,所以懂得“取舍”;正因为“会累”,所以珍惜“值得”。

(三)“双主体”何以可能:从替代逻辑到协同逻辑

在明确了二者的特质后,笔者提出“双主体”协同的理论框架。这一框架的本体论前提是:认知不等于意识,智能不等于主体性。AI 可以在认知层面扮演主体角色,而不必在意识层面取代人类。我们需要超越“主—客”二元的工具论思维,转向“主—主”协同的关系论思维——这正是我们所提出的“认知合伙人”命题的理论内核:AI 不是冷冰冰的工具,而是参与认知过程的主体性伙伴^⑧。

分布式认知理论为这一框架提供了认知科学的支撑。正如哈钦斯(Edwin Hutchins)在其开创性著作《荒野中的认知》(Cognition in the Wild)中所揭示的,人类的认知过程并非局限于颅骨之内的孤立计算,而是可以跨越个体大脑、物质工具与社会环境而分布存在^⑨。航海的领航活动并非仅存在于舵手脑中,而是分布于海图、仪器、航线与全体船员的协作之中。在这一理论视野下,AI 不再仅仅是工具,而是这一分布式系统中的新型认知节点——它与人类共同构成一个“认知共同体”。而行动者网络理论(ANT)则进一步消解了人与物的界限:拉图尔(Bruno Latour)认为,在科学知识的生产网

^⑦ 库恩关于范式转换非逻辑性的论述,参见库恩:《科学革命的结构》,金吾伦,胡新和译,北京:北京大学出版社,2003年,第12章“革命的解决”;英文版见 KUHN T S. *The Structure of Scientific Revolutions*, 3rd ed., Chicago: University of Chicago Press, 1996: 144-159. 关于“不可言说”的知识维度,参见波兰尼:《个人知识:朝向后批判哲学》,许泽民译,贵阳:贵州人民出版社,2000年;其凝练表述见 POLANYI M. *The Tacit Dimension*, Garden City, NY: Doubleday, 1966: 4.

^⑧ 同^③

^⑨ HUTCHINS E. *Cognition in the wild*, Cambridge: MIT Press, 1995: 4-5, 316, 353-354.

络中,人类与非人类同为“行动者”(actant),共同参与知识的建构^⑩。延展心智假说(Clark & Chalmers)则从哲学上论证:当外部装置稳定地参与认知过程时,它可以被视为认知系统的组成部分^⑪。三条理论脉络共同指向同一结论:认知的主体性可以在人机之间分布、延展与共享,“双主体”并非修辞,而是认知结构的实然。

“协同”在此具有三重深层含义。其一为分工性协同,各展所长——AI 负责数据处理、文献扫描与模式发现,人类负责意义诠释、价值评估与理论建构。其二为对话性协同,互为批判——AI 的高维视角可以挑战人类的思维定势,提供反直觉的“反向思考”;而人类则对 AI 的输出保持批判性审视,识别其逻辑漏洞与价值偏见。我们可以将这种关系描述为从“辅助”走向“共生”:传播者与 AI 的关系不再是“使用与被使用”,而是“协同与共生”。其三为生成性协同,共创新知——协同的产物不是二者贡献的简单相加,而是第三种知识的涌现:AI 发现数据中的隐性结构,人类赋予其理论意义,二者结合产生任何一方单独无法产出的学术洞见。

必须强调,这种双主体协同不是人类与机器的简单“分工合作”,而是两种认知逻辑的“化学反应”。在这场反应中,产物大于反应物之和。人类不再是孤独的知识探索者,而是与硅基智能共舞的编舞者。以乐队指挥为喻:过去的社会分工如同乐队里每个人拉好自己的小提琴,而今天这些具体的“演奏”技艺可能都由 AI 来完成,人的角色转变为乐队的“指挥”——指挥家不直接演奏乐器,但他决定整个乐章的走向。在这一图景中,AI 是技艺精湛的演奏家,人类是把握方向的指挥家;二者各就其位,方能奏响知识生产的交响。

二、目标的位移:学术创新中心的价值重估

人机关系的重构,必然带来对学术创新目的的重新审视。当我们承认 AI 可以成为认知主体,就必须追问:学术创新究竟为了什么?如果知识生产的参与者从“人”扩展为“人机共同体”,那么知识生产的价值坐标是否应当随之位移?本章主张,我们需要在认识论层面完成一次从“人类中心主义”向“知识本位主义”的价值位移。

(一)“人类中心主义”的知识生产范式及其限度

长期以来,学术界隐含着一种“人类中心主义”的预设:将“维护人类尊严”作为学术创新的底层动机。这种观念认为,知识生产是人类理性的特权,是人类尊严的体现。科学革命以来,“人类作为唯一理性主体”的预设深刻塑造了知识生产的制度与规范:学术署名以自然人为主体,学术荣誉授予个人,学术史书写英雄式的发现者——从哥白尼到牛顿,从达尔文到爱因斯坦,知识的发现被编织进“人类理性凯歌”的叙事之中。

这一预设并非全无道理。人类尊严确实应当得到尊重,人类理性确实值得礼赞。但问题在于:当“维护人类尊严”成为学术创新的第一性目标时,它便从一种人文关怀异化为一种认知垄断——凡是非人类生产的知识,都被预先排除在知识共同体之外;凡是挑战人类认知特权的发现,都被视为对尊严的冒犯。这种“身份歧视”在 AI 时代变得不可持续。

我们需要进行批判性追问:如果 AI 能够推导出人类从未设想、但经检验为真的理论,我们是否应因其“非人类出身”而拒绝承认其知识价值?在传播学语境下,如果 AI 从海量社交媒体数据中归纳出的受众分层模型,比传统问卷调研更精准地预测了舆论走向,拒绝它,究竟是在捍卫尊严,还是

^⑩ LATOUR B. *Science in action: how to follow scientists and engineers through society*, Cambridge: Harvard University Press, 1987: 103 - 108.

^⑪ Andy Clark and David Chalmers, “The Extended Mind”, *Analysis* 1998, 58(1): 7 - 19.

在拒绝真理？我们对此的立场是明确的：AI时代人类的角色正在从“执行处理器”升级为“战略架构师”，人类的价值不在于“亲自完成所有认知工作”，而在于“定义问题、设计系统、创造意义”。换言之，人类尊严的根基不在于“垄断认知”，而在于“驾驭认知”。把“维护人类尊严”理解为“维护人类认知的垄断权”，恰恰是对人类尊严最大的误读——尊严在于创造，而不在于排他。

（二）“知识边界突破”作为第一性目标的认识论依据

本文主张，知识边界突破应作为学术创新的第一性目标。这一主张并非反人类的“技术崇拜”，而是符合科学哲学中“真理符合论”与“实用主义真理观”的共同指向。

波普尔的“三个世界”理论为我们提供了再阐释的空间。波普尔区分了世界1（物理世界）、世界2（主观精神世界）与世界3（客观知识的世界），并强调世界3一经产生便具有自主性^⑫。理论的发现过程与其价值判断应当分离：一个理论的真值不取决于它出自谁的大脑，而取决于它与客观世界的符合程度。无论相对论是由爱因斯坦在大脑中构想，还是由AlphaFold在算法中推演，其客观真理价值是同一的。知识的“来源”不能替代知识的“效力”——正如望远镜不会因为“非人眼”而失去放大世界的资格。

默顿科学规范中的“普遍主义”（universalism）亦要求，知识判断应独立于生产者的身份特征^⑬。普遍主义是科学的精神气质的核心要素之一：科学的评价标准是“预先确定的、非个人性的标准”，而不是研究者的种族、国籍、阶级或任何身份属性。将这一规范推至AI时代：AI作为生产者不应被预先排除在知识共同体之外。判定的标准应当是：是否拓展了人类的认知边界？是否提供了新的理论解释力？是否开辟了新的研究议程？这“三个是否”构成知识本位主义的第一性判据。

在传播学研究中，AI发现的计算传播指标（如情感时间序列的拓扑特征）可能完全超越了人类的直观理解，但它可能重构我们对“舆论演化”的理解^⑭。即便人类无法直观“感受”这一发现的过程，但其知识价值不容否认。应该指出，我们正进入一个由逻辑驱动的范式窗口期——当AI能够在“可能性空间”中探索人类思维难以企及的方案组合，人类学者的任务便从“发现真理”转向“鉴定真理”与“赋予意义”。这是知识生产分工的深刻变革：AI负责“求真”的高通量探索，人类负责“求善”与“求美”的价值裁决。

因此，知识的“来源”不能替代知识的“效力”——望远镜不会因为非人眼而失去放大世界的资格，AI的发现也不会因为非人脑而失去拓展认知边疆的资格。

（三）学术创新模式的范式转换：从线性辅助到螺旋上升

基于上述价值重估，学术创新模式正经历一场深刻的范式转换。这一转换的本质在于：AI从“工具层”上升到“思维层”，人类从“执行层”回撤到“判准层”。

旧范式是线性的：人类提出问题 → 人类设计方法 → 机器执行计算 → 人类解释结果。在这一模式中，机器是被动的执行者，其角色等同于算盘与计算器——工具性的、可替换的、无主体的。这

^⑫ 波普尔：《客观知识：一个进化论的研究》，舒炜光，卓如飞，周寄中，等译。上海：上海译文出版社，2005年，第114-128，163-175页。

^⑬ 默顿：《科学社会学》，鲁旭东，林聚任，译，北京：商务印书馆，2003年，第365-377页。

^⑭ 参见周葆华，钟媛：《“春天的花开秋天的风”：社交媒体、集体悼念与延展性情感空间——以李文亮微博评论（2020—2021）为例的计算传播分析》，《国际新闻界》2021年第3期，第79-106页；张琛，马祥元，周扬，等：《基于用户情感变化的新冠疫情舆情演变分析》，《地球信息科学学报》2021年第2期，第341-350页；BAILEY M M, HEILIGMAN M I. Detecting narrative shifts through persistent structures: a topological analysis of media discourse, arXiv:2506.14836, 2025-06-14, <https://arxiv.org/abs/2506.14836>. 上述研究共同表明，AI驱动的量化方法已能揭示超越研究者直观把握的情感结构与时序特征，为舆论演化研究提供新的知识入口。

一模式与工业文明的“流水线”隐喻同构：机器在知识生产的末端提供“体力”，人类在两端提供“脑力”。

新范式则是螺旋上升的：人类确定研究方向 → AI 提出多元候选方案（包括理论假设、方法路径、解释框架） → 人类进行科学意义与伦理价值的双重判定 → 协同优化 → 螺旋迭代。在这一模式中，AI 是主动的方案提供者与认知挑战者，人类是方向的确立者与价值的裁决者。二者的关系从“指挥 - 执行”演变为“共舞 - 互构”。

在这一新模式中，“科学意义判定”成为人类学者的核心职能。这包括理论贡献度评估、逻辑一致性检验、与现实解释力的对照以及学科知识体系的适配性判断。人类不再是知识的“挖掘者”，而是新知的“鉴赏家”与“裁决者”。笔者此前关于传播者角色转型的判断在此获得学术场域的印证：传播者（研究者）的核心能力从提供“标准答案”转变为定义“关键问题”、设定 AI 探索的“价值边界”，完成从“经验执行者”到“逻辑架构师”的认知升维^⑮。笔者还进一步提出了 AI 时代人才的“四维能力坐标”：穿透问题本质的第一性原理思维（F）、批判性算法审计能力（C）、数理逻辑与系统思维的形式化基座（M）、人机转译的计算思维接口（P）——这四维能力正是“判准层”角色对研究者的新要求。

这一转换的动力学在于“螺旋上升”：每一次协同循环都产生新的认知基础，使下一轮的问题设定与方案探索在更高维度展开。例如，在未来的传播学理论创新中，可能是“AI 提出十个可能的解释模型 → 传播学者从中筛选最具理论潜力的两个 → 共同设计验证方案 → 根据结果迭代模型 → AI 再次提出反直觉的候选 → 人类重新裁决意义”。每一轮循环都不是简单重复，而是在更高的认知平面上重新开始。这就是“螺旋上升”区别于“线性循环”的根本所在：协同本身成为知识增长的内生动力。换言之，人类不再是新知的“挖掘者”，而是新知的“鉴赏家”与“裁决者”——AI 拓展认知的宽度，人类守护意义的深度。

三、责任不可让渡：价值维度的“人类锚点”

当我们欢呼于认知能力的解放与知识边界的拓展时，必须时刻警惕责任的真空。在双主体协同的架构中，确立价值维度的“人类锚点”至关重要。能力的让渡绝不意味着责任的让渡——这是双主体协同区别于“AI 取代人类”叙事的最后防线，也是“以人为本”原则在知识生产领域的制度表达。

（一）“责任主体”的哲学根基：价值对齐的不可算法化

我们的核心观点是：责任以自由意志为前提，以价值判断为内容，以社会后果的可追溯性为制度保障——这三者均无法被算法化。这一观点是责任不可让渡的哲学地基。

约纳斯的责任伦理学在当代具有极强的回响。在《责任原理》中，约纳斯指出，技术进步扩展了人类行为的时空范围，责任也随之扩大——我们不仅要当下负责，更要对遥远的未来与未知的后代负责^⑯。然而，AI 作为一个无意识、无自由意志的系统，无法成为道德责任的承担者。责任只能由具有“道德感知力”的主体承担。AI 可以成为责任的“工具”，却永远不能成为责任的“主体”：它可以被设计得符合伦理，却无法在伦理困境中“选择”伦理。

“价值对齐”（value alignment）是当前 AI 伦理的热点，但其中存在根本难题：对齐谁的价值观？如何在多元价值冲突中进行排序？这本质上是政治哲学问题，而非技术问题。在传播学语境下，如

^⑮ 同③

^⑯ 汉斯·约纳斯：《责任原理：现代技术文明伦理学的尝试》，方秋明译，世纪出版社，2013：年，第3-15，22-25页。

果 AI 主导传播政策的建议生成(如谣言判定标准、信息茧房干预策略),其背后嵌入的价值观排序——是自由优先还是秩序优先?是效率优先还是公平优先?——这绝非技术参数可以决定,而是需要在公共讨论中由人类共同体协商确定。因此,必须为 AI 装上“伦理阻尼器”,将价值导向嵌入算法设计,建立动态伦理审查机制——这一主张的深层含义正是:价值对齐的最终裁决权,必须保留在人类手中。显然,能力可以让渡,责任不可让渡——方向盘可以交给自动驾驶,但交通事故的责任书永远写着人的名字。

(二) 知识合法性的社会建构:学术共同体的“守门”功能

知识不仅是被“发现”的,更是被“承认”的。科学知识社会学(SSK)揭示,知识的有效性是在社会互动中“协商”出来的^①:一个命题要成为“知识”,必须经过学术共同体的检验、承认与制度化。同行评议制度不仅是质量把关,更是学科范式内的意义校准机制——它确保新知识能够被纳入既有的知识体系,并在对话中推动范式的演进。

当 AI 参与知识生产时,学术共同体面临着新的使命与挑战。首先,必须对 AI 生成的“伪知识”进行识别与排除。AI 极易生成看似逻辑严密实则毫无根据的“合成原创性”内容——它可能基于概率拟合拼凑出看似完美的论证链条,却缺乏真实世界的经验支撑。这种“一本正经地胡说八道”对学术共同体的辨识能力提出了更高要求:不仅检验“论证是否自治”,更要追问“证据是否真实、来源是否可靠”。我们所强调的“批判性算法审计”能力,正是学术共同体面对“合成原创性”时的第一道防线。

其次,对 AI 提出的创新方案,需要学术共同体以既有的学科规范进行“范式适配性”判断。库恩指出,科学革命是范式之间的转换,而常规科学是在范式内部的解谜^②。AI 提出的候选方案可能跨越多个范式,其理论贡献需要在既有学科框架中被定位、被检验、被承认。这一过程无法由 AI 独立完成——它需要人类学者对学科传统的深刻理解与对知识边界的敏锐感知。最后,学术共同体还需要建立新的规范来吸纳 AI 的贡献:如何在署名中体现 AI 的参与?如何在评审中评估人机协同的成果?这些规范的建立,本身就是知识合法性的社会建构过程。

如果学术界长期不加批判地采纳 AI 生成的理论框架,可能导致学科认知基础的悄然退化——一种“不知道为何相信,但相信了”的知性蒙昧。让 AI 写论文,就像让外骨骼跑马拉松——它确实更快,但人类仍然需要证明那是“自己的成绩”,并且为每一步负责。学术共同体的“守门”功能,正是为了防止这种“知性蒙昧”的发生:它确保每一项被承认的知识,都经过人类理性的检验。

(三) 法律与制度安排中的责任闭环

无论 AI 在创新过程中贡献多大,法律责任的不可转移性是现代治理体系的底线共识。法律主体资格建立在权利—义务—责任的统一性结构之上,任何缺乏自主意识与财产能力的存在者都无法嵌入其中。AI 可以成为法律关系的“客体”(如知识产权法中的“工具”),却无法成为法律关系的“主体”(如侵权责任法中的“行为人”)。

在学术场域,必须坚持可追溯性与可问责性原则。研究成果的责任人必须是可识别的自然人或法人。这要求知识产权制度做出适应与坚守。在适应层面,可考虑建立 AI 辅助成果的“贡献度披露”制度,明确区分哪些内容由 AI 生成,哪些由人类完成——正如自然科学研究中早已确立的“作者贡献声明”(author contributions statement),AI 时代的学术规范需要将“AI 贡献”纳入披露框架。在

^① 拉图尔 B, 伍尔加 S:《实验室生活:科学事实的建构过程》,刁小英,张伯霖译,北京:东方出版社,2004 年,第 209 页。

^② 库恩 T S:《科学革命的结构》,金吾伦,胡新和译,北京:北京大学出版社,2003 年,第 4 章,第 9-10 章。

坚守层面,署名权与责任主体必须是人类。2023 年以来,国际顶级期刊《自然》《科学》相继明确:大语言模型不能列为作者,AI 的使用必须充分披露^{①9}。这一全球性共识的本质,正是责任闭环的底线:署名不仅意味着荣誉,更意味着对论文内容的真实性、完整性与合法性承担法律责任——而 AI 无法被起诉、无法被问责、无法承担学术不端的后果。

传播学科在这一制度建设中可以先行先试。例如,在学术发表中建立标准的“AI 贡献声明”规范;在项目评审中明确区分“AI 方案贡献”与“人类科学判断”的比重;在学术不端调查中,始终以人类作者为责任追究终点。所以,应以“以人为本”标尺来校准人机关系——在制度层面,这意味着一切制度设计的最终目的,都是增强而非削弱人类的能动性与可问责性。技术可以改变知识生产的方式,却不能改变知识生产的责任结构:总谱永远写在人这一边。一句话,算法的尽头是合规,学术的尽头是诚实——责任闭环的最后一环,永远是人的签名。

四、传播学科的镜鉴:一个经验场域的三重启示

传播学作为研究信息流动与意义交换的学科,在这场变革中不仅提供了理论视角,更成为了人机协同创新的实验场。

(一) 研究对象的扩容:当“传播”的主体不只是人

AI 作为信息传播的新型主体,正在改写传播学的基本概念框架。从推荐算法到对话机器人,传播不再是“人与人的意义共享”,而扩展为“人—机—环境”的复杂互动系统。这意味着传播学的研究对象发生了本体论层面的扩容:算法不仅是传播的“渠道”,更是传播的“主体”;智能体不仅是内容的“载体”,更是意义的“生产者”。

这一变化要求传播学科在理论上实现突围。传统的“人类中心传播观”已无法解释 ChatGPT、DeepSeek 等大语言模型与用户的深度情感交互,也无法完全解释算法推荐对公共领域的重构。事实上,媒介的本质正从“符号通道”升级为“认知接口”,成为构成人类认知系统的必要组件;物理空间被转化为可交互的媒介界面,“环境媒介化”使空间成为“动态生成的意义生产场域”^{②0}。在这一理论视野下,传播学需要走向“广义传播生态观”,将算法、智能体纳入传播主体的范畴。在方法论上,这意味着从传统的“文本分析—受众调查”扩展为“算法审计—人机交互分析—智能体行为研究”的多元方法体系;在理论上,这意味着从“人类中心传播观”走向“人机共生传播观”——这正是“以人为本”标尺下的人机协同观在传播学对象层面的投射。概言之,当算法成为“传播者”,传播学的研究对象便从“人的意义共享”扩展为“人机环境的生态共舞”。

(二) 学科身份的再确认:传播学的“人本内核”为何不可替代

在技术喧嚣中,传播学的核心关怀——意义、权力、文化与认同——恰恰成为其不可替代的学科锚点。虽然 AI 可以模拟传播效果,生成内容,但它无法理解传播中的权力不对称。例如,算法歧视看似是代码问题,实则是社会权力结构在技术逻辑中的投射。AI 无法体认文化的深层结构,无法感受认同的情感张力,无法理解“被看见”与“被忽视”之间的政治。

传播学科在智能时代的使命,正是守护这种“意义的理解”,而非将其置换为数据的计算。在双主体协同的研究模式中,AI 可以帮助分析数以亿计的舆情数据,但如何解读这些数据背后的社会心态、文化变迁与权力博弈,依然需要传播学家的智慧。我们将“意义生产的稀缺性”列为人类在 AI 时

^{①9} 量子位:《新规:用 ChatGPT 写论文可以,列为作者不行》,百度百家号 2023-01-28 18:05 <https://baijiahao.baidu.com/s?for=pc&id=1756245018319071719&wfr=spider>

^{②0} 喻国明:《空间智能:传播学学科创新的四大维度》,《新闻与传播评论》2025 年第 4 期。

代的关键底牌——因为 AI 能处理信息,但不能生产意义;因此,传播学者的未来角色是“意义架构师”与“连接编织者”——我们不仅在塑造传播空间,更在定义人类的感知模式与文明的价值叙事。这种“人本内核”不仅是传播学科的立身之本,更是其在人机协同时代不可被算法替代的“底牌”。

应该指出传播学未来深耕的两大方向:认知神经传播学与计算传播学——前者用脑科学、心理学手段研究“用户如何接收、理解信息”,后者用大数据、算法模型研究“信息传播的规律”。这两大方向恰好构成“双主体协同”在传播学内部的学科化表达:计算传播学释放 AI 的认知优势,认知神经传播学守护人类的具身经验。传播学科的“人本内核”,正是在这种“技术之眼”与“人文之心”的协同中得以巩固与升华。显然,传播学的“人本内核”,不是对技术的抗拒,而是对意义的守护——在数据洪流中,它仍是那枚不沉的锚。

(三)新范式的雏形:传播学科如何先行探索“双主体”制度框架

传播学因其研究对象与技术的高度嵌合,最适合率先探索人机协同的制度框架。这既是学科发展的内在需求,也是其作为“经验场域”为其他学科提供范本的责任。具体而言,至少需要三个维度的制度探索:

人才培养维度:如何在研究生培养中建立“人机协同研究能力”的评估标准,不再仅仅考核传统的文献综述能力,而是考核驾驭 AI 进行理论发现的能力,以及批判性审视 AI 结果的能力。事实上, AI 时代最稀缺的人才不再是熟练的“执行者”,而是具备产品思维、系统思维、数据思维与伦理思维的“越域者”——传播学科的人才培养应当率先向“越域者”转型,锻造既懂人文又懂算法的复合型学术人才。

学术发表维度:如何在学术发表中确立 AI 贡献的规范披露机制,确保学术诚信的透明度,建立行业标准。从“作者贡献声明”到“AI 使用披露”,从“AI 生成内容的标识”到“人机协同成果的署名规范”,传播学期刊可以率先试点,为全学科提供可操作的制度样本。

科研立项维度:如何在科研立项中既鼓励 AI 辅助创新,又坚守责任主体的伦理底线。评审机制需要回答:AI 贡献如何计入创新评价?人机协同成果的原发性如何认定?责任主体如何确认?这些问题的答案,将塑造未来学术生产的制度环境。必须强调“以人为本”标尺下的人机协同——在制度层面,这意味着一切制度设计的最终目的,都是增强而非削弱人类的能动性与其可问责性。

传播学因其研究对象与技术的高度嵌合,注定成为人机协同制度探索的先行者——它的今天,是其他学科的明天。因此,传播学的探索,将为其他学科提供宝贵的经验范本,成为人机协同知识生产的先行者。从发展的角度看,传播学科正处在“范式窗口期”——谁能率先完成认知重构与制度创新,谁就能赢得定义下一个时代的权利。

结语:分工、让度与坚守——迈向负责任的人机共治

回望全文的核心逻辑,我们正站在知识生产历史性转型的关口。人工智能的深度介入,既不是人类的末日,也不是学术的终结,而是一次伟大的解放与重构。基于前文的理论推演与分析,可以在三个命题上达成共识:

第一,在能力维度上, AI 是平行的行动者。否认 AI 的认知主体性,既是认识论的保守,也是对知识突破机遇的浪费。硅基智能在数据处理、模式发现与跨域迁移上的优势,使学术创新从“人类缓慢推进”转向“双主体加速共进”。“认知合伙人”命题将获得学术场域发展创新的越来越充分的印证: AI 不再是工具,而是参与认知过程的主体性伙伴。

第二,在创新目标上,知识边界突破优先于人类尊严维护。将学术创新的目的锚定于“拓展认知

边疆”而非“守护人类垄断”，是以科学理性取代身份焦虑的正确姿态。传播学科的智能转型已经证明：AI 的深度参与并未消解传播研究的价值，反而迫使它更清醒地认识自身的“人本内核”——当计算能完成数据的搬运，意义的生产便成为人类学者最珍贵的护城河。

第三，在责任归属上，人类必须是不可让渡的锚点。这并非出于对人类优越性的幻想，而是基于价值判断不可算法化、知识合法性依赖社会建构、法律问责需要可追溯主体的三重硬约束。我们以“河流与堤坝”为喻：技术是河流，人文是堤坝——河流没有堤坝会泛滥，堤坝没有河流会干涸。在知识生产领域，AI 是奔涌的河流，人类的价值判断与责任担当是坚固的堤坝；二者的共生，才是知识生产健康发展的生态。

最终，学术创新的人机共治图景是：让 AI 释放算力之“能”，让人守住意义之“核”。唯有二者各就其位、各尽其能、各守其界，知识生产才能在智能时代实现真正的跃升——既拓展认知的边疆，又永葆人文的温度。碳基与硅基的共舞，不是谁领舞谁伴舞的主从叙事，而是两种智慧在认知舞台上各展其长、共创新章的复调交响。在这场交响中，AI 奏出的是数据的赋格，人类唱响的是意义的咏叹——而整部乐章的总谱，始终写在人这一边。