

doi:10.19920/j.cnki.jmsc.2026.07.001

# 人工智能赋能管理科学研究：范式要素转变与前沿课题<sup>①</sup>

陈国青<sup>1</sup>, 曾大军<sup>2</sup>, 王月君<sup>3</sup>, 郭迅华<sup>1</sup>, 张佳音<sup>1\*</sup>

(1. 清华大学经济管理学院, 北京 100084; 2. 中国科学院自动化研究所, 北京 100190;  
3. 北京科技大学经济管理学院, 北京 100083)

**摘要:** 人工智能正在深刻影响着管理科学研究. 本研究立足广义管理科学, 构建“范式要素 – 研究过程 – 前沿课题”框架, 阐释研究视域、研究对象、研究预设和研究方法的范式要素转变, 梳理人工智能在观测、理论构建、假说生成、模拟实验与检验等环节中的赋能路径, 以及相应用机制和边界条件. 进而, 讨论人机交互偏差循环、基于大模型的决策推演以及其他重要议题的未来探索创新方向. 研究强调, 人工智能赋能管理科学研究是在数据、任务、理论、可解释性和外部效度约束下形成的人机协同式范式演进.

**关键词:** 人工智能; 管理科学; 研究范式; 生成式大模型; 决策智能

**中图分类号:** C93      **文献标识码:** A      **文章编号:** 1007-9807(2026)07-0001-18

## 0 引言

自二十一世纪以来, 大数据等信息科技持续重塑经济运行、组织治理和科学研究方式<sup>[1-7]</sup>. 随着大模型、智能体系统和多模态学习等技术快速发展, 人工智能(AI)正在成为影响知识生产、组织决策与社会系统运行的新型通用智能基础设施. 我国2024年《政府工作报告》提出开展“人工智能+”行动, 持续推动人工智能与产业发展、科技创新和社会治理深度融合<sup>[8]</sup>; 2025年8月国务院发布《关于深入实施“人工智能+”行动的意见》<sup>[9]</sup>, 提出推动人工智能加速科学发现进程, 并探索人机协同的哲学社会科学研究方法; 2026年《政府工作报告》进一步提出“打造智能经济新形态”, 强调深化拓展“人工智能+”<sup>[10]</sup>. 人工智能已从单点技术应用进入系统性赋能阶段, 并对经

济社会运行和科研创新产生深层影响.

管理科学是理解复杂经济社会系统运行规律的重要学科. 与自然科学中相对稳定的物理对象不同, 管理科学面对的是由人、组织、制度、技术和环境共同构成的开放复杂系统, 其研究过程同时包含事实观测、行为解释、机制识别、模型推演和决策干预等多重任务. 人工智能的引入, 正在改变管理科学研究的问题识别方式、研究对象边界、模型建构逻辑和知识验证路径<sup>[11]</sup>: 一方面, 人工智能能够从海量多源数据中提取复杂语义、构造高维变量、辅助理论归纳和假说生成, 拓展管理科学对复杂现象的观测与表征能力<sup>[12-14]</sup>. 另一方面, 人工智能本身也逐渐成为管理情境中的行动者、交互者和制度性力量, 使管理科学研究对象从传统的人、组织和市场, 进一步扩展到人机共生系统、算法行为和智能决策网络.

① 收稿日期: 2026-04-06; 修订日期: 2026-05-06.

基金项目: 国家自然科学基金资助项目(72421001; 72272086; 72401027; 72293561).

通讯作者: 张佳音(1983—), 女, 陕西西安人, 博士, 长聘副教授, 博士生导师. Email: zhangjy5@sem.tsinghua.edu.cn

然而,人工智能赋能管理科学研究并不意味着对既有研究范式的简单替代。人工智能虽然在文本理解、知识生成、推理模拟和智能交互方面展现出新的能力,但其输出仍受到训练数据、模型结构、提示方式、任务设定和评价标准的共同约束,并可能带来幻觉生成、偏差放大、可解释性不足、外部效度有限以及伦理风险等问题。因此,面对人工智能的快速发展,管理科学研究当前亟需回答的关键问题是:人工智能究竟改变了管理科学哪些研究范式的哪些基本构件?这些改变如何嵌入管理科学研究过程?在何种条件下,人工智能能够形成有效、可靠且可验证的研究赋能?

已有研究从不同层面对上述问题提供了重要基础。AI for Science 相关研究强调人工智能在科学发现、实验设计、知识综合和研究自动化中的作用<sup>[11, 15-19]</sup>;社会科学研究领域开始关注人工智能对调查实验、文本分析、行为模拟和理论构建的影响<sup>[13, 20, 21]</sup>;管理科学相关领域也逐步探索大模型在变量构造、偏好识别、决策建模和管理实验中的应用<sup>[22-25]</sup>。这些研究表明,人工智能正在进入科研过程的多个环节。但总体来看,现有文献仍存在三方面不足:第一,较多研究聚焦具体工具或单一任务,对人工智能如何作用于研究范式的基本要素缺乏系统解释;第二,相关讨论多借鉴一般科学研究或社会科学研究框架,对管理科学兼具行为解释、组织情境、经济机制、优化决策和系统构建等学科特征的关注不足;第三,现有研究往往强调人工智能的赋能潜力,对相关适用条件、失效边界、验证机制和外部效度约束讨论不够。由此,有必要在已有研究基础上,从研究范式要素出发,系统分析人工智能赋能管理科学研究的内在机制和边界条件。

鉴于此,本研究立足广义管理科学范畴,构建“范式要素-研究过程-前沿课题”的分析框架,以解释人工智能驱动管理科学研究范式演进的基本逻辑。首先,围绕研究视域、研究对象、研究预设和研究方法四个范式要素,分析人工智能介入后要素转变及其作用机制,强调这种转变是在数据、

模型、理论和情境约束下形成的人机协同式范式演进。其次,从研究过程视角,梳理人工智能在现象观测、经验概括、理论构建、假说生成、模拟实验和假说检验等环节中的赋能路径,并讨论对质性建构、行为实证、计量分析、解析推演和系统构建等主流模型形态的影响。最后,聚焦人工智能引发的新议题,讨论人机交互偏差循环、基于大模型的决策推演两个重要议题以及我们新近的探索,最后展望管理科学研究的其他前沿议题和创新方向。本研究旨在为人工智能赋能管理科学研究及其创新提供前瞻性思想洞察与路径启示。

## 1 管理科学研究的内涵、要素与人工智能驱动的重构

### 1.1 广义管理科学学科与研究范式界定

管理科学是综合运用定性分析、数量建模、行为实验、计量推断、系统仿真和优化决策等方法,揭示经济管理活动规律并服务于实践的交叉性学科。狭义上,管理科学通常侧重运筹优化、决策分析、系统工程和数量模型等研究传统;广义上,管理科学则涵盖管理科学与工程、工商管理、公共管理、经济科学等相关领域,关注人、组织、市场、制度、技术与环境交互作用下的复杂经济管理活动。本研究采用广义管理科学视角,并不意在弱化不同学科分支之间的方法传统和问题差异,而是旨在从更一般的层面讨论人工智能对管理科学研究范式的共同影响及其边界条件。

范式(paradigm)通常是指特定学术共同体在一定时期内相对共享的问题意识、理论前提、概念框架、方法规范和评价标准<sup>[26, 27]</sup>。研究范式是识别研究问题、界定研究对象、设定理论前提、选择研究方法并形成可验证知识的一套相对稳定的知识生产框架。结合管理学、经济学、社会科学方法论等相关文献<sup>[28-34]</sup>,本研究将管理科学研究范式基本要素概括为:研究视域、研究对象、研究预设和研究方法,分别体现“从哪里看问题”,“研究什么”,“基于什么前提”,“如何获得和验证知识”。

这四个要素为理解人工智能条件下管理科学研究范式的变化提供一个可操作的概念框架. 其合理性在于: 任何管理科学研究, 无论采取质性建构、行为实证、计量分析、解析推演还是系统构建, 均需要在一定视域下界定研究对象, 在特定预设下开展建模与推断, 并通过相应方法形成知识主张.

1.2 管理科学研究范式的四个基本要素

1.2.1 研究视域

研究视域是指观察、解释和界定管理问题时所采用的问题边界、知识来源、情境范围和分析层次. 它回答“从哪里看问题”. 管理科学研究视域既可体现为个体、团队、组织、平台、产业、区域和国家等不同分析层次<sup>[29]</sup>, 也可体现为行为、制度、技术、文化、资源和环境等不同解释角度<sup>[35]</sup>. 研究视域决定了哪些现象被纳入研究、哪些变量被视为关键因素、哪些情境被认为具有解释意义, 也影响研究结论的适用范围和外部效度.

1.2.2 研究对象

研究对象是指被刻画、解释、预测或干预的主体、行为、关系和系统. 它回答“研究什么”. 管理科学研究对象既包括个体、团队、组织、企业、平台、市场、政府和社会系统, 也包括资源配置、组织协同、市场竞争、公共治理、风险传播和决策优化等过程<sup>[36]</sup>. 与自然科学中相对稳定的物理对象不同, 管理科学研究对象通常具有主体能动性、情境嵌入性和制度依赖性, 其行为与结果受到心理认知、组织结构、市场机制、技术环境和制度安排等

多重因素影响. 这里, 研究对象不仅包括“是什么”的静态结构特征, 也包括“做什么”的动态行为特征.

1.2.3 研究预设

研究预设是指在建模、解释、推断和验证过程中, 对主体特征、行为机制、环境条件和变量关系所作出的简化性设定<sup>[37]</sup>. 它回答“基于什么前提”. 管理科学研究预设既可体现为理性人假设、有限理性假设、效用函数、偏好结构、信息条件、分布形式和约束条件<sup>[38, 39]</sup>, 也可体现为组织目标、制度背景、行为规则、模型结构和情境边界<sup>[40, 41]</sup>. 研究预设是为了降低复杂性、明确理论边界、构建可分析模型和形成可检验命题而作出的必要抽象. 预设的合理性直接影响研究结论的解释力、稳健性和适用边界.

1.2.4 研究方法

研究方法是指获取数据、构造变量、建立模型、开展推断和验证结论的路径、工具和技术体系. 它回答“如何获得和验证知识”. 管理科学研究方法具有定性与定量、归纳与演绎、解释与预测、优化与仿真等多重特征. 典型方法形态包括质性建构、行为实证、计量分析、解析推演、优化建模、系统构建、自然实验和受控实验等. 研究方法不仅决定研究过程的技术路线, 也涵盖证据形态、理论建构方式、推断逻辑和结论可验证性<sup>[33]</sup>.

表 1 概括了四个范式要素的核心问题和基本内涵.

表 1 管理科学研究范式要素的核心问题与基本内涵

Table 1 Core questions and connotations of Management Sciences research paradigm elements

| 范式要素 | 核心问题      | 基本内涵                   |
|------|-----------|------------------------|
| 研究视域 | 从哪里看问题    | 问题边界; 知识来源; 情境范围; 分析层次 |
| 研究对象 | 研究什么      | 主体; 行为; 关系; 系统         |
| 研究预设 | 基于什么前提    | 主体特征; 行为机制; 环境条件; 变量关系 |
| 研究方法 | 如何获得和验证知识 | 数据获取; 变量构造; 模型建构; 推断验证 |

1.3 人工智能驱动的范式要素转变

人工智能发展正在改变着管理科学研究的问题识别、对象建模、前提设定和方法体系. 人工智能驱动的范式要素转变是在特定数据条件、任务结构、理论边界和验证机制下对于范式要素不同

程度的重组和演进. 本研究将这种转变概括如下.

1.3.1 研究视域转变：从领域内分析到跨域扩展

人工智能对研究视域的影响, 主要体现为从相对封闭的领域内分析, 逐步走向多源数据、跨域

知识和高维语义空间的融合分析. 人工智能的引入能够从文本、图像、语音、平台行为、社会网络、宏观环境和制度文本等多源异构数据中提取信息, 并将原本难以表征的语义、情绪、认知和情境因素纳入分析框架.

具体说来, 第一, 人工智能增强了跨数据源整合能力, 能够引入更多外部变量和情境信息. 例如, 大模型可通过语义抽取和逻辑推理, 辅助识别社会情绪、管理者意图与资本市场波动之间的关联<sup>[42]</sup>. 第二, 人工智能拓展了变量和构念的表征空间, 使低维、显性变量能够扩展到高维、隐性和动态变量空间. 例如, 深度学习方法可用于提取宏观经济指标与微观企业特征之间的复杂交互<sup>[43]</sup>; 多模态表征可用于刻画领导力风格、员工心理状态等潜在特征<sup>[44]</sup>. 第三, 跨域扩展也带来新的挑战, 例如变量解释性下降、跨域知识错配、算法黑箱、数据偏差和价值对齐问题等, 其有效性受到数据质量、理论约束和验证机制的影响.

#### 1.3.2 研究对象转变: 从人类主体到人机共生

人工智能对研究对象的影响, 主要体现为从人、组织、市场和制度等传统主体, 进一步扩展到人机共生系统、算法行为和智能决策网络. 这里, 广义管理科学并非始终以人类行为为唯一核心. 在运筹优化、系统工程和决策分析等传统中, 资源配置、系统效率和优化方案同样是重要研究对象. 人工智能带来的关键变化是机器不再只是被动工具或技术背景, 而是在越来越多的管理情境中成为能够生成内容、影响行为、参与互动和决策的行动性主体因素.

具体说来, 第一, AI 系统本身成为新的研究对象. 例如, 数字平台中的智能算法、自动驾驶系统中的路径规划、大模型的内容生产、智能代理的自主决策等, 均构成需解释和评估的新型对象<sup>[45-47]</sup>. 第二, 人类主体与 AI 系统的交互关系成为新的研究对象. 人工智能被嵌入组织管理、平台治理、市场交易和公共决策过程后, 会影响个体认知、组织协同、劳动控制、消费者偏好和社会信任. 例如, 算法管理可能影响从业者的技术压力、心理契约和工作满意度<sup>[48]</sup>, AI 系统的广泛应用也会引发公平性、透明性和责任归属等治理问题<sup>[49]</sup>. 因此, 这一转变不仅扩大了研究对象范围, 也要求

重新审视主体性、能动性、责任归属和价值对齐等基础问题.

#### 1.3.3 研究预设转变: 从显性结构预设到复合预设体系

人工智能对于研究预设的影响, 主要体现为预设形态的再结构化. 人工智能的引入, 在部分场景中降低了对显性预设的依赖, 但同时也引入了新的隐性预设, 包括数据来源预设、训练分布预设、模型架构预设、提示词预设、评价函数预设和平台环境预设.

具体说来, 人工智能推动管理科学研究从主要依赖显性结构预设, 转向显性预设与隐性预设并存的复合预设体系. 一方面, 相关方法可在一定条件下辅助评估异质性处理效应<sup>[50]</sup>, 刻画消费者偏好变化<sup>[51, 52]</sup>, 并在“预测-优化”框架下支持复杂资源配置问题<sup>[53]</sup>. 另一方面, 这些结果仍依赖样本覆盖、数据生成机制、模型训练过程、任务定义和评价指标. 如果隐性预设与研究情境不匹配, 模型输出可能产生偏差、误判或不可解释结论. 由此可见, 人工智能的引入将部分预设从显性结构转移到数据、模型和计算过程之中.

#### 1.3.4 研究方法转变: 从工具辅助到智能驱动

人工智能对研究方法的影响, 主要体现为从单点工具辅助逐步走向智能驱动的研究链条重组.

具体说来, 第一, 在数据处理和变量构造环节, 人工智能可用于大规模文本、图像和行为数据的标注、分类、聚类 and 语义萃取. 例如, 大模型可对消费者评论、新闻报道和社交媒体内容进行情感识别、主题分类和构念测量<sup>[13]</sup>. 第二, 在模型构建和行为模拟环节, 人工智能可用于生成个体或群体行为代理, 模拟消费者偏好、组织互动、平台竞争和市场反馈等过程. 例如, 大模型可模拟消费者对产品属性的感知差异与动态态度, 生成消费者态度与购买意向分布<sup>[20]</sup>. 第三, 在推断检验和实验设计环节, 人工智能可与结构因果模型、优化算法和仿真平台结合, 用于辅助构建虚拟实验场、生成反事实情境和检验管理假说<sup>[21]</sup>. 值得一提的是, 智能驱动需通过人工复核、模型稳健性检验、跨数据集验证和外部效度评估等方式保证研究可靠性.

1.4 范式要素转变的作用机制与边界条件

综上,人工智能对管理科学研究范式要素的影响可以概括为“扩展-重构-迁移-协同”的作用机制,如表 2 所示.此外,表 2 概括了要素转变及其适用范围.其中,边界条件既包括 AI 赋能生效的前提,也包括影响赋能效果强度的调节因素,二者共同构成范式要素转变的约束空间.在研究视域上,跨域扩展的有效性取决于数据质量,当多源异构数据存在系统性偏差或语义不一致时,跨域整合可能放大偏差;同时,知识匹配和理论约束则决定扩展后的变量能否转化为有效理论洞察.在研究对象上,主体边界的界定是人机共生分

析的前提,AI 在具体管理情境中被界定为工具性辅助、代理性行动者或是制度性力量,直接影响因果归因的逻辑;此外,责任归属和价值对齐作为调节因素,影响研究结论规范性和可信度.在研究预设上,样本覆盖构成复合预设有效性的前提,训练数据代表性不足则隐性预设可能偏离研究情境;可解释性和情境适配则决定复合预设能否被有效识别、审查和调整.在研究方法上,稳健性是智能驱动的前提,模型输出对提示词、参数或版本高度敏感时,结论可复现性将受到挑战;外部效度和伦理规范则调节研究结论适用范围和社会可接受度.

表 2 人工智能驱动的管理科学研究范式要素转变

Table 2 AI-driven shifts in paradigm elements of Management Sciences research

| 范式要素 | 传统形态   | AI 驱动转变 | 作用机制                               | 边界条件           |
|------|--------|---------|------------------------------------|----------------|
| 研究视域 | 领域内分析  | 跨域扩展    | 人工智能扩展了管理科学可观测、可表征和可解释的变量空间        | 数据质量;知识匹配;理论约束 |
| 研究对象 | 人类主体   | 人机共生    | 人工智能与人类共同成为管理情境中的行动性主体因素           | 主体边界;责任归属;价值对齐 |
| 研究预设 | 显性结构预设 | 复合预设体系  | 人工智能推动预设从显性理论结构向数据、模型和计算过程的复合体系迁移  | 样本覆盖;可解释性;情境适配 |
| 研究方法 | 工具辅助   | 智能驱动    | 人工智能促进研究链条从人工主导的分段式流程向智能驱动的迭代式流程演进 | 稳健性;外部效度;伦理规范  |

2 人工智能赋能管理科学研究的过程路径与效度边界

本节从管理科学研究过程出发,分析人工智能如何在不同环节形成赋能路径,并进一步讨论其适用条件、潜在风险、验证要求与效度边界.

2.1 管理科学研究过程

研究过程是研究范式在时间维度和操作层面的展开.人工智能对范式要素的影响,最终会体现为对研究过程的重组.一般而言,管理科学研究过程可以概括为从“现象观测”到“理论归纳”的归纳过程,以及从“理论命题”到“经验检验”的演绎过程,两者共同构成螺旋式上升的知识生产逻辑(如图 1)<sup>[54, 55]</sup>.具体说来,研究者首先通过测量、抽样、编码、统计或案例分析等方式对现象进行观测,并在样本概括、变量构造和参数估计基础上形成经验概括;随后,通过概念凝练、机制解释和命题构造,将经验概括提升为理论;进一步,理论在

特定情境下通过逻辑演绎形成可检验假说;最后,研究者通过实验、调查、计量分析、仿真推演或优化求解等方式对假说进行检验,并根据结果修正经验认识和理论框架.

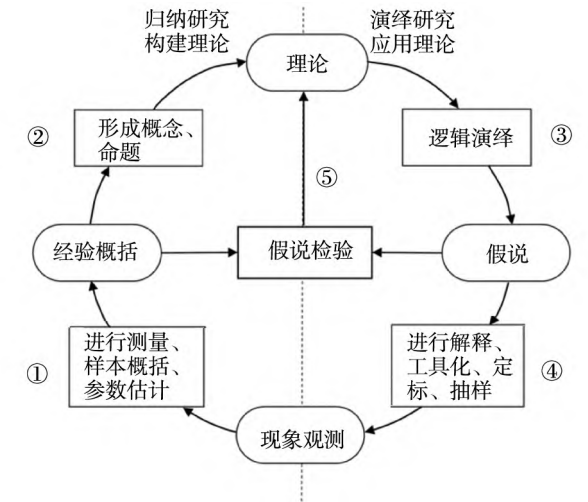


图 1 管理科学研究过程

Fig. 1 Research process of Management Sciences

注：图中椭圆代表信息主体或研究对象,方框代表动作或研究行动,箭头代表过程转换;后续图形含义相同.

## 2.2 人工智能赋能管理科学研究过程的五类路径

人工智能赋能管理科学研究过程体现为知识生产过程中关键转换环节的重塑,对应着图 1 的五类路径:

### 2.2.1 从现象观测到经验概括:智能语义萃取与构念测量

从现象观测到经验概括,关键在于将复杂管理现象转化为可分析的数据、变量和构念.传统研究中非结构化数据的编码、标注和解释往往高度依赖人工判断,存在成本高、规模受限、一致性不

足和跨情境迁移困难等问题.大模型(如 LLM)的引入能够对文本、图像、语音、评论、新闻、访谈和平台行为等多源数据进行语义解析、主题识别、情感分类和构念测量.相关研究从文本标注、语义度量、跨领域标注和高阶语义识别等方面验证了大模型在语义萃取中的应用潜力<sup>[13, 56-58]</sup>.这里,大模型不仅替代部分人工编码工作,更重要的是扩展了对非结构化现象的测量能力,使原本难以进入模型的语义、情绪、态度、认知和制度等因素转化为可分析变量.其一般路径可概括为图 2(即对应图 1 中路径①).

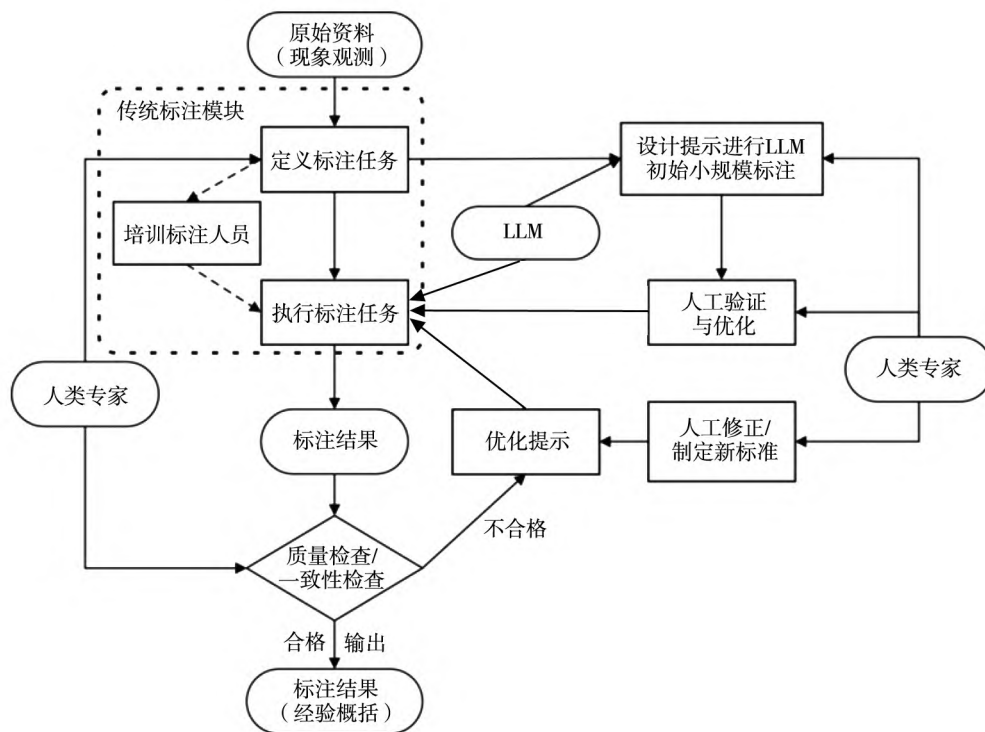


图2 大模型智能化语义萃取

Fig. 2 LLM-enabled semantic extraction

### 2.2.2 从经验概括到理论构建:概念归纳与命题生成

从经验概括到理论构建,关键在于从经验材料中提炼概念、识别关系并形成命题.传统理论建构依赖研究者的理论敏感性、经验积累和归纳能力,具有较强的解释深度,但也面临材料规模受限、编码过程难以复现和跨领域知识整合不足等问题.人工智能可通过大规模文本分析、概念聚

类、语义比较和跨文献综合,发现潜在模式、提取候选概念并生成初步理论命题.已有研究从概念归纳、归纳推理、知识综合和认知基础模型等方面,探索人工智能辅助理论构建的可能性<sup>[59-62]</sup>.这意味着,人工智能可以在理论构建中发挥“候选概念生成器”和“跨域知识连接器”的作用,帮助扩展理论搜索空间.其一般路径可概括为图 3(即对应图 1 中路径②).

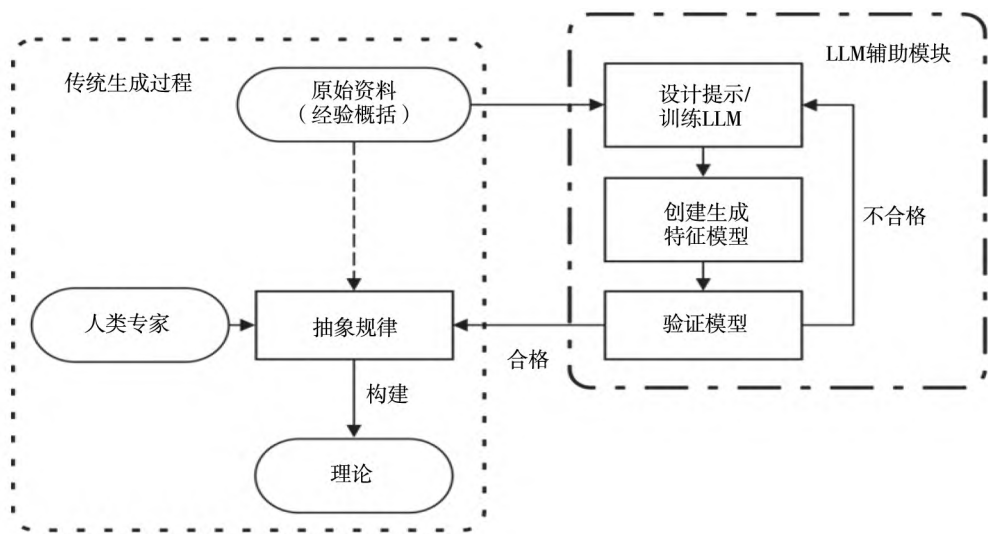


图 3 大模型辅助理论构建  
Fig. 3 LLM-assisted theory building

2.2.3 从理论内生到假说生成：增强型逻辑演绎  
与建模求解

从理论到假说生成,关键在于基于既有理论框架,对特定情境下的变量关系、因果机制或行为结果进行逻辑演绎.大模型的推理、检索和生成能力,使其能够在文献综合、机制推断、变量关联识别和研究假说表达中发挥作用.已有研究尝试结合大模型与因果知识图谱生成潜在假说,或通过多渠道知识获取、显式推理链和“生成-辩论-进化”机制辅助研究假说与实验设计生成<sup>[63-67]</sup>.

在管理科学中,理论到假说的演绎并不只表现为文字命题生成,还包括优化模型、博弈模型、决策模型和仿真模型的构建.对于运筹优化和博弈分析等研究,大模型可辅助完成从自然语言情境到数学表达、从模型结构到代码实现、从求解结果到解释说明的转化.例如,已有研究开始探索利用大模型诊断不可行优化问题、改进自动化优化建模能力、通过强化学习和知识对齐提升优化求解表现,以及辅助推导博弈均衡和形式化证明<sup>[22-25]</sup>.其一般路径可概括为图 4(即对应图 1 中路径③).

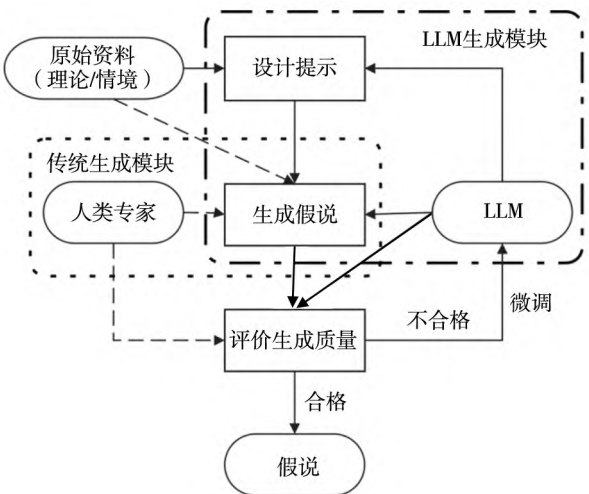


图 4 大模型增强型逻辑演绎  
Fig. 4 LLM-enhanced logical deduction

2.2.4 从假说生成到现象观测：智能体模拟与研究设计

从假说生成到现象观测,关键在于将理论假说转化为可观察、可测量和可检验的研究设计,包括情境解释、变量操作化、实验设计、样本构造和测量工具开发等.人工智能可通过构建虚拟代理和多智能体系统,模拟个体、组织或群体在特定情境中的反应,为实验设计、情境推演和样本扩展提供新路径.已有研究利用大模型构建生成式代理,模拟个体反应、实验结果、心理能力、主观幸福感以及博弈互动



### 2.3 人工智能赋能路径的适用条件与验证要求

人工智能介入研究过程多个环节,其有效赋能取决于任务类型、数据条件、理论基础、验证机制和应用情境等。以下阐释不同环节的赋能路径、适用条件、潜在风险和验证要求。

从现象观测到经验概括,人工智能主要通过智能语义萃取辅助构念识别、文本标注和变量测量<sup>[13, 56, 57, 74]</sup>。其适用性存在差异,效果取决于概念语义边界清晰度、语料质量可靠性、标注规范可操作性,人工校验样本<sup>[13, 56-58]</sup>等情形。对于模糊、争议性强或依赖情境理解的构念,可通过设计人机协作标注分析、领域理论指导、专家审查及外部效度检验<sup>[57, 58]</sup>,提升大模型标注与构念测量的可靠性和可解释性。

从经验概括到理论构建,人工智能可辅助发现经验模式、整合文献证据和生成候选概念关系,适用于研究范围和分析任务较清晰的场景,如大样本经验模式识别、文献综合及标准化认知实验行为预测<sup>[59-62]</sup>。缺乏连贯上下文信息或跨领域知识整合时,可能产生貌似合理但缺乏机制深度的“伪理论”。而模式解释力、适用范围及其与其他理论的区分度等判断仍依赖研究者分析与学术共同体评价<sup>[60, 75]</sup>。

从理论到假说生成,人工智能可辅助生成候选假说、扩展机制推演和形成可检验的模型表达。可通过结构化知识、推理链、问题情境、经验数据或形式化模型等方法优化并约束大模型生成过程<sup>[25, 63-67]</sup>。若研究问题无理论覆盖、变量关系和作用机制不确定,或对生成的假说缺乏系统审查,可能生成表面化、逻辑不连贯、难以操作化或检验的假说<sup>[64-67]</sup>。研究者可设计因果图谱推演、结构化推理链、多智能体辩论进化等增强性推演过程,通过理论判断、约束诊断和生成验证循环等评估验证<sup>[22, 25, 63, 66, 67]</sup>,并依据可检验性和伦理边界筛选和修正假说。

从假说生成到现象观测,人工智能通过智能体模拟辅助研究设计、情境推演和机制预演。在个

体信息、实验情境边界、行为规则较清晰且结果可校验的场景中,模拟易发挥可靠作用<sup>[68, 69, 71]</sup>。但许多复现实验本质上仍是在验证大模型能力,不等同于真实场景再现。受训练语料、提示设定和内部机制影响,模型输出可能表现出多样性差、提示敏感、刻板印象、文化偏差或推理不稳定等问题<sup>[14, 67-71, 76]</sup>。需与真实实验结果、历史调查数据或专家判断对齐校准。

在假说检验中,人工智能适用于研究问题操作化明确、易于建模记录、干预机制较清晰、结果变量可观测比较的场景,如依托学科领域实验构建AI被试并模拟检验、构建反事实机器人以创造比较情境,或作为真实信息实验中的干预工具<sup>[21, 72, 73]</sup>。需强调的是,自动化实验验证不能保证有效的因果识别,模拟数据无天然外部效度,也不能自动解决选择偏差、混杂变量、样本代表性和测量误差等问题。涉及真实管理决策、消费者行为、组织互动和公共政策效果的研究,需通过随机实验、自然实验、准实验设计、真实行为数据或多方法交叉验证来增强结论信度。

人工智能介入研究过程,不仅改变研究工具,更重塑研究者角色定位、研究思维和认知习惯。传统研究中研究者兼顾设计与执行,AI介入后研究者重心转向人机协作的组织设计、逻辑引导与质量把控。此外,可靠知识的形成需接受构念效度、内部效度、外部效度、可复现性以及伦理与治理规范等多维评价,从构念测量、因果解释、跨情境适用、可重复验证和伦理规范等方面综合判断AI输出的有效性。例如,对于标注分类,应重点关注构念效度和一致性;对于假说生成和理论构建,应重点关注理论解释力、可证伪性和内部效度;对于智能体模拟和自动化实验,应重点关注外部效度、可校准性和跨情境稳定性;对于涉及真实组织、用户或公共决策的研究,还需特别关注数据合规、伦理约束和责任归属。

### 2.4 对管理科学主流模型形态的影响

上述赋能路径进一步影响管理科学主流模型

形态的演化. 这里, 主流模型形态并非指单一具体模型, 而是指管理科学研究中长期形成的若干知识生产方式. 以下将讨论五类典型模型形态及其演化和创新方向.

#### 2.4.1 质性建构模型: 从人工编码到人机协同的语义归纳

在质性研究中, 大模型不仅是编码工具, 更是智能驱动的理论共建者. 它将研究视域从有限文本片段拓展至海量多模态数据的高阶语义空间, 能够识别出人类研究者难以察觉的隐性模式. 研究方法上, AI 作为协作主体参与编码与概念提炼, 实现了人机协同的深度归纳. 这也使得研究预设降低对先验理论的依赖, 可从访谈语料中提取核心概念并发现其内在关系, 为理论构建提供有效支撑. 总体而言, 其演化方向是形成以研究者解释为中心、以大模型语义归纳为辅助的人机协同理论建构机制.

#### 2.4.2 行为实证模型: 从静态实验到人机互动情境中的行为模拟

在研究视域上, AI 通过构建人机共生模拟环境, 实现了对真实管理情境的跨域仿真. 在研究对象上, 大模型不仅能基于智能驱动生成大规模异质性样本, 模拟目标群体的心理感知与决策行为, 还能自动设计复杂实验场景. 同时, 这一过程将研究预设的显性个体理性程度/同质性偏好转为隐性构念关系, 通过反事实推断与生成式实验等研究方法, 检验复杂的因果假说, 突破传统实验室的外部效度瓶颈. 总体而言, 其演化方向是将智能体模拟作为实验设计和机制探索的前置工具.

#### 2.4.3 计量分析模型: 从结构化变量到多模态构念重构

在研究视域上, 计量模型在 AI 赋能下, 实现了从低维结构化数据向高维非结构化数据的跨域跃迁. 在研究对象上, AI 能够从多模态语料中提取情绪、态度等高阶潜变量, 将原本难以量化的社会心理特征纳入计量体系. 这打破了传统模型对变量形式的严格研究预设, 允许引入复杂动态的

新型变量关系. 在研究方法上, AI 可实现从现象观测到经验概括的智能转化, 使得对复杂管理现象的解释力显著提升, 以揭示传统方法无法捕捉的深层机制. 总体而言, 其演化方向是从单一结构化数据建模走向多模态构念重构与机制识别相结合.

#### 2.4.4 解析推演模型: 从人工建模到大模型辅助的形式化表达与求解

在研究预设上, AI 能够在很大程度上摆脱对传统预设的依赖, 辅助构建博弈论或运筹优化模型. 这也将研究对象扩展至包含 AI 算法本身的复杂人机系统, 通过大规模仿真模拟进行模型验证. 在研究方法上, AI 使得模型构建不再受限于研究者的数学表征能力边界, 能够处理更高维度、更具非线性特征的复杂系统优化问题, 拓展了管理科学解析模型的应用边界. 总体而言, 其演化方向是形成“大模型理解与表达 - 专业求解器计算 - 研究者理论审查”的协同机制.

#### 2.4.5 系统构建模型: 从信息系统开发到智能化管理系统创生

AI 通过数据层增强分析与建模层语义理解, 将研究视域从现实空间拓展至新兴管理场景的虚实融合空间, 并帮助快速定位理论缺口, 实现质性归纳与行为编程的智能化. 在研究对象方面, 大模型构建在推理层与优化层上具有自主决策能力的人机共生智能体, 通过多智能体交互与深度强化学习, 自动演绎系统动态演化过程. 这不仅柔化了研究预设, 更通过算子智能协同对接专业求解器, 实现从“解释世界”到“智能改造世界”的研究方法闭环. 总体而言, 其演化方向是实现可解释、可审计、可验证和可治理的智能管理系统创生.

### 3 人工智能赋能管理科学研究的前沿议题

上述讨论从范式要素和研究过程两个层面, 分析了人工智能赋能管理科学研究的一般机制.

本节进一步讨论这些机制在前沿议题中的体现。

### 3.1 人机交互偏差循环与对齐机制

在人机交互情境中,模型输出、人类行为和数据反馈之间存在着偏差循环情形<sup>[77]</sup>。其关键机制为:人工智能等机器学习方法在数据采集、模型训练和结果生成过程中,可能形成采样偏差、拟合偏差、生成偏差和评估偏差等模型偏差;模型输出的预测、推荐、生成内容或决策建议,又会影响用户

的感知、态度、行为选择和观点表达,形成选择偏差、表达偏差和认知依赖等行为偏差。一方面,行为偏差被记录为新的用户行为数据,并在后续训练中被再次学习;另一方面,模型偏差在管理情境中被持续输出,进一步影响主体行为和系统反馈(图7)。这一循环不仅影响人机交互行为与后果,也影响着研究数据的真实性、变量测量的可靠性和实验结果的外部效度。

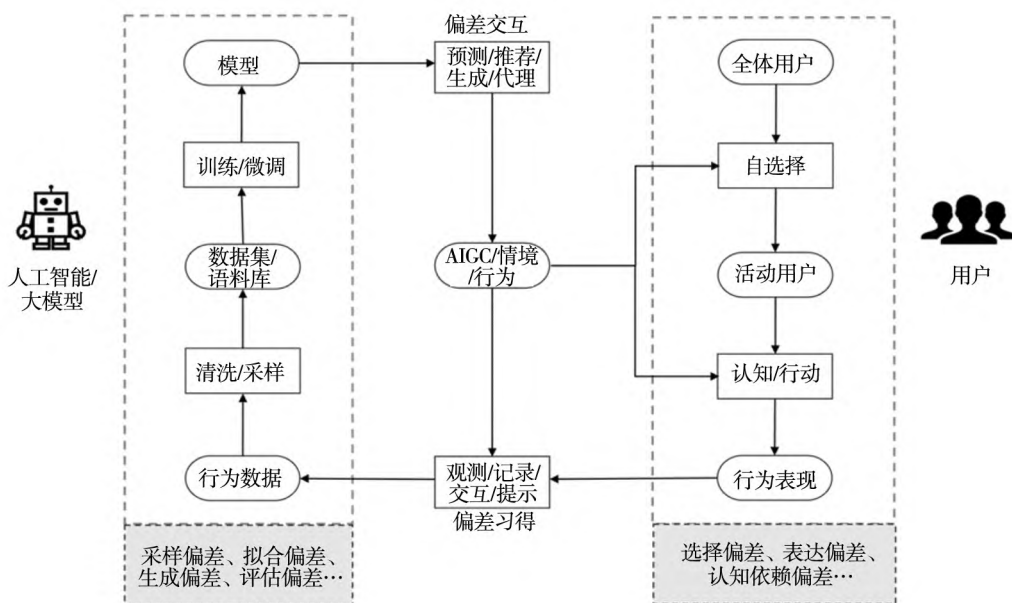


图7 人机交互偏差循环

Fig. 7 Bias iteration in human-machine interaction

面对人机交互偏差循环,解决途径不仅包括模型层面的技术修正,还包括在要素层面构建对齐机制;在研究视域上,需要识别跨域数据、生成内容和外部语料引入的偏差来源;在研究对象上,需要区分人类主体行为、模型生成行为和人机交互行为;在研究方法上,需要通过人工复核、因果识别、跨模型验证和真实情境校准等方式,降低偏差对研究结论和管理行为的影响。结合管理科学研究场景,本研究将讨论下面两类对齐思路。

#### 3.1.1 人智认知对齐

人智认知对齐旨在保证人工智能参与研究判断时,其输出能够被研究者理解、解释、审查和修正。其核心问题在于:当大模型参与文本编码、主题识别、概念归纳和理论构建时,如何保证模型生

成结果不是表面语义相似或统计共现,而是能够被研究者解释、验证并纳入学术推理过程的知识判断。

在质性分析、案例研究和文本研究等场景中,可以构建“AI生成-人工审查-模型修正-再生成”的人在回路机制(图8)。具体而言,研究者首先依据理论框架和研究问题定义概念边界、编码规则和评价标准;随后,LLM对文本材料进行初步编码、主题提取和解释生成;研究者对模型输出进行一致性、准确性和理论解释力评估,并据此修正提示词、编码规则或概念定义;最后,通过多轮迭代形成相对稳定的编码体系和理论解释结果。该过程重点是通过人机协同提高质性研究的透明性、一致性和可复核性。

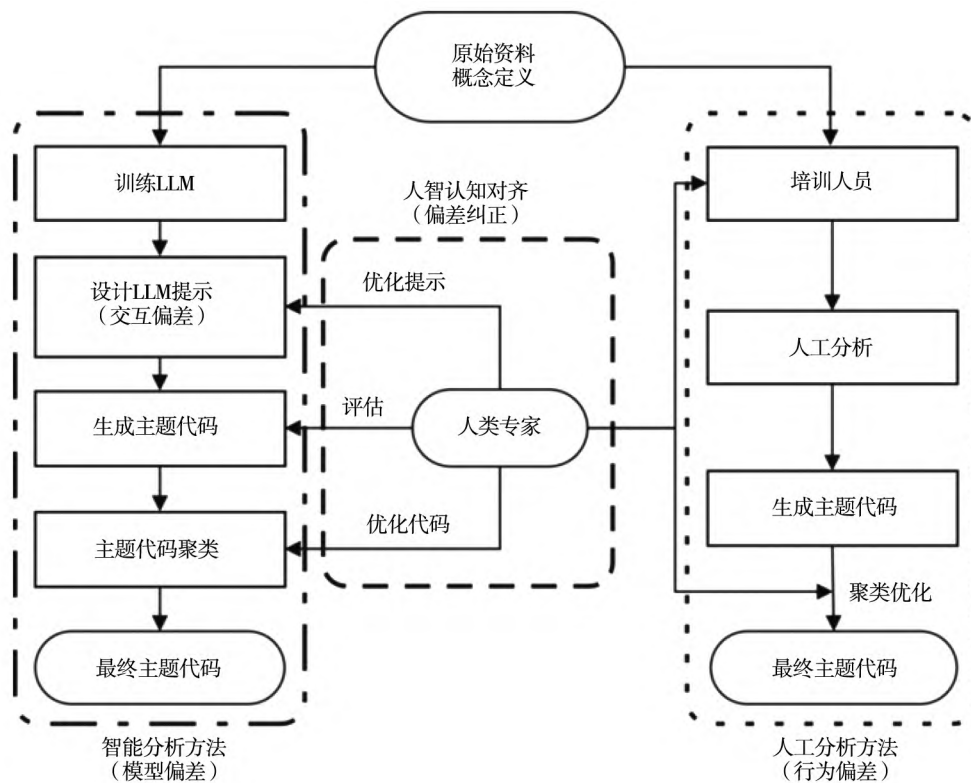


图 8 人智认知对齐 (以主题分析为例)

Fig. 8 Human-AI cognitive alignment: An example of thematic analysis

### 3.1.2 虚实情境行为对齐

虚实情境行为对齐旨在保证人工智能模拟的主体行为、互动过程和情境反应能够与真实管理情境保持可校准的一致性. 其核心问题在于: 当大模型被用于构建研究代理、模拟用户行为或推演群体互动时, 如何避免智能体行为脱离真实情境, 进而导致研究结论偏误.

在相关研究中, 可以采用“大模型 + 领域模型 + 真实数据校准”的方法 (图 9). 一方面, LLM 用于生成角色画像、行为反应和互动过程, 扩展研究者对复杂情境的模拟能力; 另一方面, 需要引入领域理论、行为规则、历史数据和专家知识, 对 LLM 生成的智能体行为进行校准和约束. 例如, 在平台创作者激励、消费者干预、组织协同或公共政策沟通等场景中, LLM 可以模拟不同主体的认知状态、动机结构和行为反应, 但其结果需与真实平台数据、调查数据、实验结果或专家判断进行对齐. 模拟结果经过校准和验证, 才可作为研究设计

优化、机制探索和策略预评估的辅助依据.

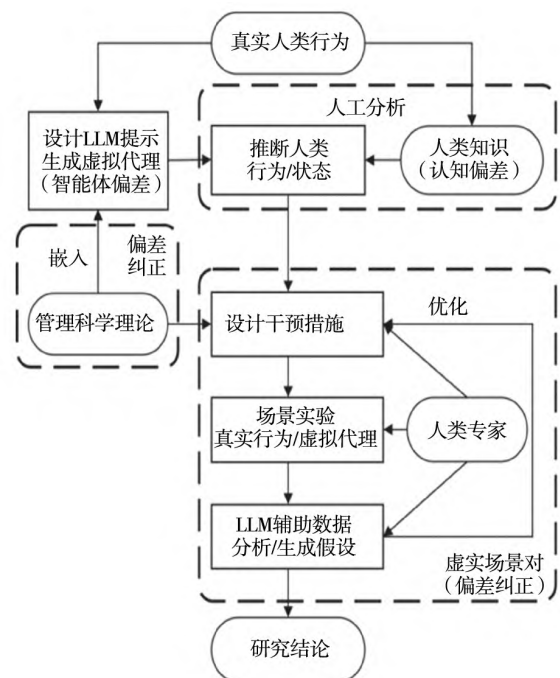


图 9 虚实情境行为对齐 (以内容创作者干预为例)

Fig. 9 Virtual-real behavioral alignment: An example of content creator intervention

### 3.2 基于大模型的决策推演

决策是管理科学的核心要义.决策推演是管理决策研究和应用的重要议题,也是人工智能赋能复杂管理问题建模、仿真和策略评估的典型场景.随着决策环境呈现高不确定性、多主体互动、跨域变量耦合和动态反馈等特征,传统推演方法在情境建模、主体刻画、策略生成和结果解释方面面临新挑战.

大模型与通用自主智能体为决策推演提供了新的方法基础.这里“通用”是指智能体构建范式在不同管理情境中的可迁移、可复用和可配置能力.区别于传统方法主要依赖预设规则和局部交互,大模型驱动的通用自主智能体能够基于角色设定、领域知识、任务目标和环境反馈,形成情境理解、目标推理、策略生成、工具调用和行为修正的闭环过程,从而更适合刻画复杂管理决策中主体认知、组织互动和动态演化.

研究价值主要体现在:第一,在研究视域上,大模型可整合政策文本、行业数据、市场信息、舆情语料和专家知识,支持复杂决策情境的跨域建模;第二,在研究对象上,通用自主智能体可构建具有角色身份、目标函数、认知状态、行为规则和反馈机制的决策主体,从而刻画人、组织、平台、政府和算法系统之间的互动关系;第三,在研究预设上,大模型可从非结构化语境中提取决策约束、行动逻辑和偏好线索,并将部分显性预设转化为数据、模型、提示词和情境设定中的隐性预设;第四,在研究方法上,大模型可与知识图谱、因果模型、优化模型、多智能体仿真和专业求解器结合,形成从情境理解、主体建模、策略生成、仿真推演到结果评估的闭环流程.

基于上述逻辑,本研究将基于大模型的决策推演概括为一种通用型大模型决策推演架构.该架构由六个基本模块构成:情境建模用于界定决策问题、目标、主体和约束;知识增强用于整合事实、规则、案例和领域理论;通用自主智能体构建用于刻画人、组织、平台、政府或算法系统的目标、

认知、能力、行为规则和反馈机制;推演规划用于生成可比较的策略路径和情景组合;仿真实验用于观察不同策略下的系统演化;策略评估用于开展稳健性分析、反事实比较和专家审查.

特别地,基于大模型的推演结果作为决策支持和研究证据的一部分,不能替代决策主体的责任判断、价值权衡和制度审查.此外,对于复杂管理决策问题,推演结果通常需接受三类检验:一是事实一致性检验,即模型输入和情境设定是否与真实数据和领域知识一致;二是机制合理性检验,即智能体行为规则和因果链条是否符合理论逻辑;三是结果稳健性检验,即推演结论是否在不同参数、场景和模型设定下保持稳定.

### 3.3 其它未来探索创新方向

除了人机交互偏差循环和基于大模型的决策推演议题之外,从更广泛意义上讲,人工智能迅猛发展催生了一系列亟待探索的新议题与广阔创新空间,包括下列若干未来探索方向.

#### 3.3.1 管理科学理论发展与构建

人工智能的融合渗透正在深入管理科学的理论根基,催生面向“人机共生”的基础理论创新,亦意味着在研究视域、预设、方法上不同程度的转变.关键科学问题是经典管理科学理论体系(如基于“理性人”假设的决策理论、以组织边界为前提的战略理论、以人类代理为中心的领导理论)如何回应 AI 成为“准行动者”乃至“协同主体”的现实,进而实现理论拓展、修正或重构.重要研究内容包括:AI 介入下经典理论前提(如有限理性、代理理论)的重新检验与边界拓展;AI 智能体承担管理职能的行为模式特征及人机协同行为的微观理论构建;围绕智能经济新形态,构建能够解释人机复杂互动、智能涌现以及新价值源泉的理论框架.

#### 3.3.2 复杂管理决策建模方法创新

大模型、因果推断、多智能体仿真、优化求解和知识图谱等方法的结合,为复杂管理决策建模提供了新的方法体系,也是范式要素转变在方法

论层面的综合体现. 关键科学问题是刻画不同决策主体的角色身份、认知状态、目标偏好、行为规则和反馈调整机制; 表征主体之间的关系网络、制度约束、资源结构、事件演化和情境变量; 生成可比较的策略路径、仿真方案和评估指标. 重要研究内容包括: 自然语言情境到形式化模型的可靠转换; 多源异构数据到决策变量的构念测量; 多主体互动过程中的机制识别; 大模型与专业求解器之间的协同.

### 3.3.3 数据 - 模型 - 内容治理

人工智能带来的数据、模型算法和生成内容, 正成为管理科学研究和应用创新的关键资源, 也体现了对范式要素的全方位影响. 相应地, 治理问题也从传统数据治理扩展到模型算法治理和内容治理. 特别地, 诸如“算法歧视”、大模型“数据投毒”(如生成式引擎优化 GEO) 引发的认知操纵等管理伦理问题已成为治理的关键焦点. 需深入探讨 AI 向善的实现路径与伦理边界, 将负责任 AI 原则融入 ESG (环境、社会和治理) 战略, 以防范技术滥用对公平和伦理的侵蚀. 关键科学问题是在保障算法创新与数据流动效率的同时, 构建覆盖数据确权、算法透明度及内容合规性的多维协同治理机制. 重要研究内容包括: 训练数据和生成内容的权属与责任界定; 大模型决策中的偏差、歧视和幻觉的识别与纠正; 模型开发者、部署者、使用者和决策者之间的责任边界划分; 面对数据投毒、深度伪造、GEO 等外部对抗风险, 敏捷治理和风险预警机制构建.

### 3.3.4 智能管理系统与价值创造

智能管理系统(如: 具身智能机器人、人机混合团队、智能客服、智能供应链和智能公共服务

等) 的发展正重塑人与机器之间的协作方式, 集中体现了范式要素转变在虚实空间的协同影响. 关键科学问题是探究人工智能改变组织流程、产业协同、服务模式和价值创造机制. 重要研究内容包括: 具身机器人等具有物理实体和类人特征的智能体如何通过“临场感”影响人类心理认知、团队认同与协作; 人机混合团队如何改变人类员工的自我效能感并引发领导力结构与沟通模式的适应性变革; 智能服务中的新型信任关系特征、行为机理以及价值分配机制.

## 4 结束语

人工智能赋能管理科学研究正由工具增强走向范式演进. 本研究立足广义管理科学范畴, 构建“范式要素 - 研究过程 - 前沿课题”的分析框架, 系统讨论人工智能赋能管理科学研究的内在机制、实现路径与边界条件. 在范式要素层面, 人工智能推动研究视域从领域内分析走向跨域扩展, 推动研究对象从传统主体扩展到人机共生, 推动研究预设从显性结构预设转向显性与隐性预设并存, 推动研究方法从单点工具辅助走向人机协同的智能驱动. 在研究过程层面, 人工智能通过智能语义萃取、概念归纳、逻辑演绎、智能体模拟和实验验证等路径赋能管理科学知识生产过程, 同时构念效度、内部效度、外部效度、可复现性以及伦理与治理有效性等多维检验的重要性凸显. 最后, 从前沿议题视角讨论了人机交互偏差循环、基于大模型的决策推演以及其它未来研究探索创新方向. 本研究旨在为人工智能赋能管理科学研究及其创新提供前瞻性思想洞察与路径启示.

## 参 考 文 献:

- [1] Frankel F, Reid R. Big data: Distilling meaning from data[J]. Nature, 2008, 455(7209): 30.
- [2] Hilbert M, López P. The world's technological capacity to store, communicate, and compute information[J]. Science, 2011, 332(6025): 60 - 65.
- [3] Science Staff. Challenges and opportunities[J]. Science, 2011, 331(6018): 692 - 693.
- [4] 陈国青. 未来信息化发展方向[N]. 人民日报(理论版), 2013 - 01 - 09.

- Chen Guoqing. The Development Direction of Informatization in the Future [N]. People's Daily (Theoretical Edition), 2013-01-09. (in Chinese)
- [5] 陈国青, 吴刚, 顾远东, 等. 管理决策情境下大数据驱动的研究和应用挑战——范式转变与研究方向[J]. 管理科学学报, 2018, 21(7): 1-10.
- Chen Guoqing, Wu Gang, Gu Yuandong, et al. The challenges for big data driven research and applications in the context of managerial decision-making: Paradigm shift and research directions[J]. Journal of Management Sciences in China, 2018, 21(7): 1-10. (in Chinese)
- [6] 陈国青, 张维, 任之光, 等. 面向大数据管理决策研究的全景式 PAGE 框架[J]. 管理科学学报, 2023, 26(5): 4-22.
- Chen Guoqing, Zhang Wei, Ren Zhiguang, et al. Panoramic PAGE framework for big data research on managerial decision-making[J]. Journal of Management Sciences in China, 2023, 26(5): 4-22. (in Chinese)
- [7] 陈国青, 任明, 卫强, 等. 数智赋能: 信息系统研究的新跃迁[J]. 管理世界, 2022, 1: 180-195.
- Chen Guoqing, Ren Ming, Wei Qiang, et al. Data-intelligence empowerment: A new leap of information systems research[J]. Journal of Management World, 2022, 1: 180-195. (in Chinese)
- [8] 国务院. 2024 年政府工作报告[R]. 北京: 国务院, 2024.
- State Council. Report on the Work of the Government 2024[R]. Beijing: State Council, 2024. (in Chinese)
- [9] 国务院. 国务院关于深入实施“人工智能+”行动的意见[R]. 北京: 国务院, 2025.
- State Council. Opinions of the State Council on Further Implementing the “AI+” Initiative[R]. Beijing: State Council, 2025. (in Chinese)
- [10] 国务院. 2026 年政府工作报告[R]. 北京: 国务院, 2026.
- State Council. Report on the Work of the Government 2026[R]. Beijing: State Council, 2026. (in Chinese)
- [11] Grossmann I, Feinberg M, Parker D C, et al. AI and the transformation of social science research[J]. Science, 2023, 380(6650): 1108-1109.
- [12] Heidt A. AI for research: The ultimate guide to choosing the right tool[J]. Nature, 2025, 640(8058): 555-557.
- [13] Gilardi F, Alizadeh M, Kubli M. ChatGPT outperforms crowd workers for text-annotation tasks[J]. Proceedings of the National Academy of Sciences, 2023, 120(30): e2305016120.
- [14] Kosinski M. Evaluating large language models in theory of mind tasks[J]. Proceedings of the National Academy of Sciences, 2024, 121(45): e2405460121.
- [15] Farrell H, Gopnik A, Shalizi C, et al. Large AI models are cultural and social technologies[J]. Science, 2025, 387(6739): 1153-1156.
- [16] Bail C A. Can generative AI improve social science? [J]. Proceedings of the National Academy of Sciences, 2024, 121(21): e2314021121.
- [17] Messeri L, Crockett M J. Artificial intelligence and illusions of understanding in scientific research[J]. Nature, 2024, 627(8002): 49-58.
- [18] Wang H, Fu T, Du Y, et al. Scientific discovery in the age of artificial intelligence[J]. Nature, 2023, 620(7972): 47-60.
- [19] Xu Y, Liu X, Cao X, et al. Artificial intelligence: A powerful paradigm for scientific research[J]. The Innovation, 2021, 2(4): 100179.
- [20] Estevez M, Ballestar M T, Sainz J. Market research and knowledge using Generative AI: The power of Large Language Models[J]. Journal of Innovation & Knowledge, 2025, 10(5): 100796.
- [21] Manning B S, Zhu K, Horton J J. Automated social science: Language models as scientist and subjects[J]. arXiv Preprint arXiv: 2404.11794, 2024. DOI: 10.48550/arXiv.2404.11794.
- [22] Chen H, Constante-Flores G E, Li C. Diagnosing infeasible optimization problems using large language models[J]. INFORM: Information Systems and Operational Research, 2024, 62(4): 573-587.
- [23] Huang C, Tang Z, Hu S, et al. ORLM: A customizable framework in training large models for automated optimization modeling[J]. Operations Research, 2025, 73(6): 2986-3009.
- [24] Xie W, Hu Y X, Xu J, et al. MURKA: Multi-reward Reinforcement Learning with Knowledge Alignment for Optimization Tasks[C]//The Thirty-ninth Annual Conference on Neural Information Processing Systems, 2025.

- [25] Liu W, Zhou X, Wang X, et al. An LLM agent-based framework for analytical characterization of Nash equilibria[J]. Nexus, 2025, 3(1): 100107.
- [26] Kuhn T S, Hacking I. The Structure of Scientific Revolutions[M]. Chicago: University of Chicago Press, 1970.
- [27] 陈国青, 曾大军, 卫 强, 等. 大数据环境下的决策范式转变与使能创新[J]. 管理世界, 2022, (2): 95 – 105.  
Chen Guoqing, Zeng Dajun, Wei Qiang, et al. Transitions of decision-making paradigms and enabled innovations in the context of big data [J]. Journal of Management World, 2022, (2): 95 – 105. (in Chinese)
- [28] Ansoff H I. The emerging paradigm of strategic behavior[J]. Strategic Management Journal, 1987, 8(6): 501 – 515.
- [29] 罗 珉. 管理学范式理论述评[J]. 外国经济与管理, 2006, 28(6): 1 – 10.  
Luo Min. A review of paradigm theory in management[J]. Foreign Economics & Management, 2006, 28(6): 1 – 10. (in Chinese)
- [30] 田国强. 现代经济学的基本分析框架与研究方法[J]. 经济研究, 2005, (2): 113 – 124.  
Tian Guoqiang. The basic analytical framework and methodologies in modern economics[J]. Economic Research Journal, 2005, (2): 113 – 124. (in Chinese)
- [31] 王海文. 范式的演进与现代经济学的发展[J]. 经济评论, 2006, (5): 104 – 107.  
Wang Haiwen. Paradigm evolution and the development of modern economics[J]. Economic Review, 2006, (5): 104 – 107. (in Chinese)
- [32] 李怀祖. 管理研究方法论[M]. 西安: 西安交通大学出版社, 2000.  
Li Huaizu. Methodology of Management Research[M]. Xi'an: Xi'an Jiaotong University Press, 2000. (in Chinese)
- [33] 袁 方. 社会研究方法教程[M]. 北京: 北京大学出版社, 1997.  
Yuan Fang. A Course in Social Research Methods[M]. Beijing: Peking University Press, 1997. (in Chinese)
- [34] 杨 耕. 社会科学方法的发生、范式及其历史性转换[J]. 中国社会科学, 1994, (1): 135 – 145.  
Yang Geng. The emergence, paradigm, and historical transformation of social science methodology[J]. Social Sciences in China, 1994, (1): 135 – 145. (in Chinese)
- [35] 波特伦·谢弗德. 马克思、桑巴特、韦伯与现代资本主义起源争论[J]. 经济社会史评论, 2023, (4): 4 – 12.  
Bertram Schefold. Marx, Sombart, Weber and the debate about the genesis of modern capitalism[J]. Economic and Social History Review, 2023, (4): 4 – 12. (in Chinese)
- [36] 刘炯忠. 论管理科学研究的对象、特点与主导方法[J]. 教学与研究, 1989, (3): 41.  
Liu Jiongzhong. The object, characteristics and dominant methods of management science research[J]. Teaching and Research, 1989, (3): 41. (in Chinese)
- [37] Nkwake A M. Why are assumptions important? [M]//Working with Assumptions in International Development Program Evaluation: With a Foreword by Michael Bamberger. New York: Springer New York, 2012: 93 – 111.
- [38] 罗 珉. 论管理学范式革命[J]. 当代经济管理, 2005, 27(5): 35 – 40.  
Luo Min. The revolution of paradigms in management[J]. Contemporary Economic Management, 2005, 27(5): 35 – 40. (in Chinese)
- [39] 陈荣秋, 周水银. 生产运作管理的理论与实践[M]. 北京: 中国人民大学出版社, 2002.  
Chen Rongqiu, Zhou Shuiyin. Theory and Practice of Production and Operations Management[M]. Beijing: China Renmin University Press, 2002. (in Chinese)
- [40] 宋冬林. 从范式危机看经济学的发展[J]. 当代经济研究, 1997, (2): 5 – 9.  
Song Donglin. On the development of economics from the perspective of paradigm crisis[J]. Contemporary Economic Research, 1997, (2): 5 – 9. (in Chinese)
- [41] 顾明毅. 论幸福对微观经济学效用最大化法则的修正[J]. 经济问题, 2009, (9): 20 – 23.  
Gu Mingyi. The microeconomic adjustment of hedonistic utility maximization[J]. On Economic Problems, 2009, (9): 20 – 23. (in Chinese)
- [42] de Kok T. ChatGPT for textual analysis? How to use generative LLMs in accounting research[J]. Management Science, 2025, 71(9): 7888 – 7906.
- [43] Chen L, Pelger M, Zhu J. Deep learning in asset pricing[J]. Management Science, 2024, 70(2): 714 – 750.
- [44] Cheng S Y, Golshan N M. Silent suffering: Using machine learning to measure CEO depression[J]. Journal of Accounting Research, 2025, 63(2): 689 – 767.

- [45] Wang C, Shi Y, Guo X, et al. Probing digital footprints and reaching for inherent preferences: A cause-disentanglement approach to personalized recommendations[J]. *Information Systems Research*, 2025, 36(3): 1314–1332.
- [46] Hou J, Wang L, Wang G, et al. The double-edged roles of generative AI in the creative process: Experiments on design work[J]. *Information Systems Research*, 2025. DOI: 10.1287/isre.2024.0937.
- [47] Walsh M, Preus A, Gronski E. Does ChatGPT have a poetic style? [J]. *arXiv preprint arXiv: 2410.15299*, 2024.
- [48] Cram W A, Wiener M, Tarafdar M, et al. Examining the impact of algorithmic control on Uber drivers' technostress[J]. *Journal of Management Information Systems*, 2022, 39(2): 426–453.
- [49] Hu X, Huang Y, Li B, et al. Human-algorithmic bias: Source, evolution, and impact[J]. *Management Science*, 2026, 72(1): 495–514.
- [50] Chernozhukov V, Demirer M, Duflo E, et al. Reply to: Comments on “Fisher-Schultz Lecture: Generic machine learning inference on heterogeneous treatment effects in randomized experiments, with an application to immunization in India” [J]. *Econometrica*, 2025, 93(4): 1177–1181.
- [51] Goli A, Singh A. Frontiers: Can large language models capture human preferences? [J]. *Marketing Science*, 2024, 43(4): 709–722.
- [52] Wu J, Liu H, Yao X, et al. Unveiling consumer preferences: A two-stage deep learning approach to enhance accuracy in multi-channel retail sales forecasting[J]. *Expert Systems with Applications*, 2024, 257: 125066.
- [53] Elmachtoub A N, Grigas P. Smart “predict, then optimize” [J]. *Management Science*, 2022, 68(1): 9–26.
- [54] Wallace W. *The Logic of Science in Sociology* [M]. London: Routledge, 1971.
- [55] 陈晓萍, 徐淑英, 樊景立. 组织与管理研究的实证方法 [M]. 第2版, 北京: 北京大学出版社, 2012.  
Chen Xiaoping, Xu Shuying, Fan Jingli. *Empirical Methods in Organizational and Management Research?* [M]. 2nd ed., Beijing: Peking University Press, 2012. (in Chinese)
- [56] Le Mens G, Kovács B, Hannan M T, et al. Uncovering the semantics of concepts using GPT-4[J]. *Proceedings of the National Academy of Sciences*, 2023, 120(49): e2309350120.
- [57] Carlson N A, Burbano V. The use of LLMs to annotate data in management research: Foundational guidelines and warnings[J]. *Strategic Management Journal*, 2026, 47(3): 699–725.
- [58] Liu R, Li J, Zivkovic M, et al. Automating in high-expertise, low-label environments: Evidence-based medicine by expert-augmented few-shot learning[J]. *MIS Quarterly*, 2025, 49(3): 1049–1094.
- [59] Lam M S, Teoh J, Landay J A, et al. Concept induction: Analyzing unstructured text with high-level concepts using lloom [C]//*Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*. 2024: 1–28.
- [60] Shrestha Y R, He V F, Puranam P, et al. Algorithm supported induction for building theory: How can we use prediction models to theorize? [J]. *Organization Science*, 2021, 32(3): 856–880.
- [61] Asai A, He J, Shao R, et al. Synthesizing scientific literature with retrieval-augmented language models[J]. *Nature*, 2026. DOI: 10.1038/s41586-025-10072-4.
- [62] Binz M, Akata E, Bethge M, et al. A foundation model to predict and capture human cognition[J]. *Nature*, 2025, 644(8078): 1002–1009.
- [63] Tong S, Mao K, Huang Z, et al. Automating psychological hypothesis generation with AI: When large language models meet causal graph[J]. *Humanities and Social Sciences Communications*, 2024, 11(1): 1–14.
- [64] Wang Q, Downey D, Ji H, et al. Scimon: Scientific inspiration machines optimized for novelty [C]//*Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*. 2024: 279–299.
- [65] Liu H, Zhou Y, Li M, et al. Literature meets data: A synergistic approach to hypothesis generation [C]//*Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*. 2025: 245–281.
- [66] O'Neill C, Ghosal T, Răileanu R, et al. Sparks of science: Hypothesis generation using structured paper data[J]. *arXiv preprint arXiv: 2504.12976*, 2025.
- [67] Gottweis J, Weng W H, Daryin A, et al. Towards an AI co-scientist[J]. *arXiv preprint arXiv: 2502.18864*, 2025.
- [68] Park J S, Zou C Q, Shaw A, et al. Generative agent simulations of 1 000 people[J]. *arXiv preprint arXiv: 2411.10109*, 2024.
- [69] Hewitt L, Ashokkumar A, Ghezae I, et al. Predicting results of social science experiments using large language models[J]. <https://samim.io/dl/Predicting%20results%20of%20social%20science%20experiments%20using%20large%20language%20models>

- 20models. pdf, 2024.
- [70] Pataranutaporn P, Powdthavee N, Archiwaranguprok C, et al. Simulating human well-being with large language models: Systematic validation and misestimation across 64 000 individuals from 64 countries[J]. *Proceedings of the National Academy of Sciences*, 2025, 122(48): e2519394122.
- [71] Akata E, Schulz L, Coda-Forno J, et al. Playing repeated games with large language models[J]. *Nature Human Behaviour*, 2025, 9(7): 1380 – 1390.
- [72] Hosseinmardi H, Ghasemian A, Rivera-Lanas M, et al. Causally estimating the effect of YouTube’s recommender system using counterfactual bots[J]. *Proceedings of the National Academy of Sciences*, 2024, 121(8): e2313377121.
- [73] Velez Y R, Green D P, Sevi S. Chatbot Voting Advice Applications inform but seldom sway young unaligned voters[J]. *Proceedings of the National Academy of Sciences*, 2025, 122(50): e2515516122.
- [74] Wright A G C, Ringwald W R, Vize C E, et al. Assessing personality using zero-shot generative AI scoring of brief open-ended text[J]. *Nature Human Behaviour*, 2026: 1 – 15.
- [75] Lu C, Lu C, Lange R T, et al. Towards end-to-end automation of AI research[J]. *Nature*, 2026, 651(8107): 914 – 919.
- [76] Zewail A, Figueroa A, Graham J, et al. Moral stereotyping in large language models[J]. *Proceedings of the National Academy of Sciences*, 2026, 123(10): e2519941123.
- [77] 郭迅华, 吴 鼎, 卫 强, 等. 机器学习与用户行为中的偏差问题: 知偏识正的洞察[J]. *管理世界*, 2023, 39(5): 145 – 159.
- Guo Xunhua, Wu Ding, Wei Qiang, et al. The problem of biases in machine learning and user behavior: Insights into knowing the deviated and upholding the undeviated[J]. *Journal of Management World*, 2023, 39(5): 145 – 159. (in Chinese)

## Artificial Intelligence for Management Sciences: Paradigm-element shifts and emerging issues

CHEN Guo-qing<sup>1</sup>, ZENG Da-jun<sup>2</sup>, WANG Yue-jun<sup>3</sup>, GUO Xun-hua<sup>1</sup>, ZHANG Jia-yin<sup>1\*</sup>

1. School of Economics and Management, Tsinghua University, Beijing 100084, China;

2. Institute of Automation, Chinese Academy of Sciences, Beijing 100190, China;

3. School of Economics and Management, University of Science and Technology Beijing, Beijing 100083, China

**Abstract:** Artificial Intelligence is profoundly affecting Management Sciences research. From a broader perspective of Management Sciences, this paper develops a “paradigm elements-research process-emerging issues” framework, explains the shifts in light of research scopes, research objects, research assumptions, and research methods, and examines AI-enabled pathways across observation, theory building, hypothesis generation, simulation experiments, and hypothesis testing, along with related mechanisms and boundary conditions. Finally, it discusses key research subjects including human-computer interaction biases and large model-based decision simulation, as well as other directions of future exploration. This paper argues that AI empowering Management Sciences research represents a human-AI collaborative paradigm evolution constrained by data conditions, task structures, theoretical boundaries, interpretability, and external validity.

**Key words:** Artificial Intelligence; Management Sciences; research paradigm; generative large models; decision intelligence