

人工智能驱动的社会科学有趣性： 理论“有趣性”的综述与启示

□ 周贤永 刘玉欣

领域编辑推荐语：

文章对理论“有趣性”文献进行了系统的梳理，所提出的人类与人工智能组成的知识共同体建构“有趣的集体共识”的路径，以及可能存在的风险，对我们利用人工智能提升社会科学理论“有趣性”具有启示和指导意义。

——贾良定

摘 要：随着“人工智能驱动的社会科学研究”（AI for Social Science, AI4SS）日益深化，如何借助人工智能驱动社会科学的理论“有趣性”已成为重要问题。对此，本研究系统综述理论“有趣性”并挖掘启示。研究发现，理论“有趣性”的核心在于否定受众的既定假设，但其局限于话语生产而忽视了社会建构性。理论“有趣性”是知识共同体在知识传统与制度情境下，通过话语协商达成的一种集体共识。人类与人工智能组成的知识共同体通过“扎根知识传统—适应制度情境—话语协商有趣”路径建构“有趣的集体共识”；特别地，需要警惕人工智能幻觉、算法趋同、批判惰性三大风险。本研究有助于弥合经典社会科学哲学与当代人工智能方法论之间的理论断裂，并为借助人工智能增强我国社会科学理论“有趣性”提供实践启示。

关键词：人工智能；人工智能驱动的社会科学研究；理论“有趣性”；社会科学

一、问题提出

随着人工智能（Artificial Intelligence, AI）的快速发展，“人工智能驱动的社会科学研究”（AI for Social Science, AI4SS）正在兴起（Xu et al., 2024；马亮，2025a）。2025年8月，《国务院关于深入实施“人工智能+”行动的意见》发布，

强调“推动哲学社会科学研究方法向人机协同模式转变”^①。然而，人工智能正在促成大量同质化的研究（谢耘耕和张伶俐，2025）。《自然》（*Nature*）的最新研究考察了1980—2025年4130万篇论文，发现人工智能使得研究更加同质化（Hao et al., 2026）。相比自然科学，社会科学更加依赖文本，人工智能输出相似内容的特性（Van Quaquebeke, 2025），使得其更易产生同质化研究。

过去二十年，中国社会科学研究越来越规范，但趣味性却并未同步增强（马亮，2025b）。这一现象并非中国独有，众多国际社会科学研究被批评为无趣的“学术狗屁”（Kirchherr, 2023）。在人工智能推动知识生产趋于同质化的背景下，“无趣”的社会科学可能会进一步增加。管理学具有强实践性（白长虹等，2025），并以问题为导向（盛昭瀚，2019；贾良定等，2024），其理论价值不仅取决于正确性，还依赖于是否被受众关注和应用。因此“有趣”就成为管理学顶级期刊的重要评审标准（Kilduff, 2006；Tihanyi, 2020），甚至到了过于关注的地步（陈晓萍和沈伟，2018）。社会科学的知识生产同样具有社会性（Trigg, 1985），且人工智能时代下具有保护人类价值理性的独特使命（李智超等，2026；赵乃萱和李江，2025）。因此，“有趣性”成为一个理论吸引受众、实现价值的必要条件。

1971年，Murray Davis（1971）提出理论“有趣性”，指出“有趣的理论是那些否定受众既定假设的理论”（Davis, 1971）。这一理论深刻影响了国外社会科学的理论和实践，一项调查表明，57%的管理学顶级期刊《美国管理学

学报》（*Academy of Management Journal*, AMJ）编委会成员认为，一篇文章之所以被评选为“最有趣”，是因为它否定了受众的既定假设（Bartunek et al., 2006）。在人工智能生产同质化社会科学研究背景下，Davis（1971）的理论为如何增强我国社会科学研究“有趣性”提供了重要工具。然而，国内学者对理论“有趣性”的研究却较为匮乏，已有的主要集中于应用层面。白胜（2014，2015）两次讨论了其在管理研究中的运用策略，并与严茜（2015）将其细化到了会计学的应用中，张志学等（2014）则具体化到了组织行为学的研究中。因此，国内研究尚且缺少系统性的综述。

基于此，本研究旨在综述 Davis（1971）提出的理论“有趣性”，回答如何借助人工智能驱动社会科学的理论“有趣性”。本研究有助于在理论上弥合经典社会科学哲学与当代人工智能方法论之间的断层，在实践上为社科学者借助人工智能增强理论“有趣性”提供可行路径。

二、理论“有趣性”的研究综述

（一）理论“有趣性”的主要内容

1971年，Davis（1971）发表文章“*That's Interesting! Towards a Phenomenology of Sociology and a Sociology of Phenomenology*”，揭示理论“有趣性”的核心在于否定受众的既定假设。具体内容包括以下几点：一是明确“有趣”是指能够吸引人们注意力。二是辨明有趣和无趣的

^① 中华人民共和国中央人民政府．国务院关于深入实施“人工智能+”行动的意见．2025年8月26日．https://www.gov.cn/zhengce/content/202508/content_7037861.htm，访问日期：2026年5月23日。

理论，指出有趣的理论通过挑战人们习以为常的认知框架，打破“现象”与“实在”之间的常规对应关系，让人恍然大悟并形成“真是有趣！”的感受，从而引发关注与讨论；无趣的理论则往往重复或肯定受众已有直觉，导致被斥为“显而易见”“无关紧要”或“荒谬”。三是提出“有趣性索引”（The index of the interesting），Davis（1971）归纳出12类“有趣命题”的基本形式，包括组织、构成、抽象、泛化、稳定、功能、评价、关联、共存、协变、对立

和因果，以及3种无趣的理论命题形式，并为每一种形式举出经典例子进行佐证，如表1所示。四是指出理论“有趣性”的群体差异，有趣感受高度依赖于受众的背景假设，不同社会群体（如普通大众与专业学者）的认知基础存在差异，因此同一理论可能在不同群体中引发截然不同的反应。五是以现象学为导向开辟了“有趣社会学”（Sociology of the Interesting），以补充传统知识社会学，聚焦理论接受过程中的心理转变。

表1 理论“有趣性”和“无趣性”的命题形式（Davis, 1971）

类别	基本形式	具体形式	例子	受众反应
理论 “有趣性”	组织	看似无组织的现象实际上是有组织的现象，反之亦然	斐迪南·滕尼斯在《共同体与社会》中指出，人与人之间的关系看似杂乱无章，实际上可以组织为“有机的团结”与“机械的团结”两种主要类型	真是有趣！
	构成	看似各种各样的异质现象，实际上是由单一元素构成的，反之亦然	西格蒙德·弗洛伊德指出，儿童行为、原始人行为、神经症患者行为、群体行为、梦境、笑话、口误等看似无关的现象，实际上都是由同一本能驱动的	
	抽象	看似个体现象的东西实际上是整体现象，反之亦然	埃米尔·涂尔干在《自杀论》中指出，自杀看似是个体行为，实际上是社会过程	
	泛化	看似局部现象，实际上却是普遍现象，反之亦然	卡尔·曼海姆在《意识形态与乌托邦》中指出，思想过程的意识形态限制和扭曲，看似只影响资产阶级，实际上影响所有阶级	
	稳定	看似稳定不变的现象，实际上是不稳定变化的现象，反之亦然	卡尔·马克思在《资本论》中指出，资产阶级社会的社会组织看似是永恒的，实际上即将发生剧变	
	功能	看似是一种无效的现象，实际上却是一种有效的现象，反之亦然	罗伯特·默顿在《社会理论与社会结构》中指出，政治机器看似无法有效实现社区目标，实际上却非常有效	
	评价	看似不好的现象实际上是好的现象，反之亦然	R. D. 莱因在《经验的政治》中指出，精神分裂症看似是坏事，实际上却是好事	
	关联	看似不相关（独立）的现象实际上是相关（相互依赖）的现象，反之亦然	奥古斯特·霍林斯黑德在《社会阶级与精神疾病》中指出，社会阶级与精神疾病看似无关，实际上相关	

续表

类别	基本形式	具体形式	例子	受众反应
理论 “有趣性”	共存	看似可以共同存在的现象实际上却不可能共同存在，反之亦然	德尼·德·鲁日蒙在《西方世界的爱》中指出，爱情与婚姻看似可以共存，实际上无法共存	真是有趣！
	协变	看似正相关的现象实际上是负相关的，反之亦然	大卫·卡普洛维茨在《穷人付得更多》中指出，较低收入群体在许多商品和服务的支出看似减少，实际上却增加	
	对立	看似相似（几乎相同）的现象实际上是相反的现象，反之亦然	马歇尔·麦克卢汉在《理解媒介》中指出，广播与电视看似相似，实际上是相反的媒介	
	因果	在因果关系中看似独立的现象，实际上却是因变量，反之亦然	霍华德·贝克尔在《局外人》中指出，某些个体的反常行为看似导致了别人给其贴上“越轨者”的标签，但实际上，这种行为恰恰是被人贴上标签之后才产生的	
理论 “无趣性”		看起来是这样，实际上就是这样	丈夫经常影响妻子的政治行为	这太明显了！
		事实与你一直以来的认知毫无关联	爱斯基摩人比犹太人更可能……	这无关紧要！
		所有看似真实的事实根本就不是真实的；你一直以来认为正确的所有事实实际上都是错误的	社会因素对人的行为没有影响	这太荒唐了！

（二）理论“有趣性”的发展脉络与学科影响

Davis（1971）提出理论“有趣性”，随之越来越多的国外学者参与到理论“有趣性”的讨论中。国内学者近十年来才开始关注到 Davis（1971）的“有趣性”理论，研究相对匮乏，相较之下，国外的理论发展更为丰富和连续。因此主要回顾国外的文献脉络。据 ResearchGate 的数据，截至 2025 年 12 月 6 日，引用 Davis（1971）的文章高达 1257 篇，剔除预印等重复文章，引用 Davis（1971）文章的文献 1215 篇^①。

图 1 展示了被引的变化趋势。从图 1 可以看出，理论“有趣性”的发展呈现出三阶段。1971—1999 年为理论“有趣性”发展的萌芽

期，年引用文章数量多徘徊于个位数，但 1983 年出现一次小高峰，达到 30 篇，其中具有代表性的研究是 Carr（1983）对事件结构、兴趣、重要性与连贯性的讨论，显示该议题已开始引发学术反思；此后引用数量在 1989 年为 16 篇、1999 年为 14 篇。尽管总体关注有限，但学者已初步关注反直觉问题，如 Locke 和 Golden - Biddle（1997）对“问题化”机制的探讨，Brekhus（1998）对“理所当然”社会现象的关注，以及 Davis（1999）关于概念魅力的论述。2000—2010 年是增长期，引用数量稳步上升，2005 年为 10 篇，2006 年跃升至 30 篇，并于 2010 年达到 35

^① 引用统计仅基于英文语种且收录于 ResearchGate 的期刊论文、图书，不包含以非英语语种发表或未被 ResearchGate 收录的相关研究。

篇,表明理论“有趣性”逐渐获得学界关注和认可。2011年至今进入波动期,2011年引用数量达60篇,2020年达到77篇高峰后有所回落,但始终维持在30篇以上。这一阶段,学者们通过建设和他人的关系(Dutton & Dukerich, 2006)、神秘隐喻(Alvesson & Kärreman, 2007)、

反思相反论点(Sparrowe & Mayer, 2011)、采用“问题化”策略(Alvesson & Sandberg, 2011, 2013; Hafermalz et al., 2020)、重新审视既定假设(Leatherbee et al., 2020)以及引入异质视角(Božič, 2020)等,不断丰富理论“有趣性”的实现路径。

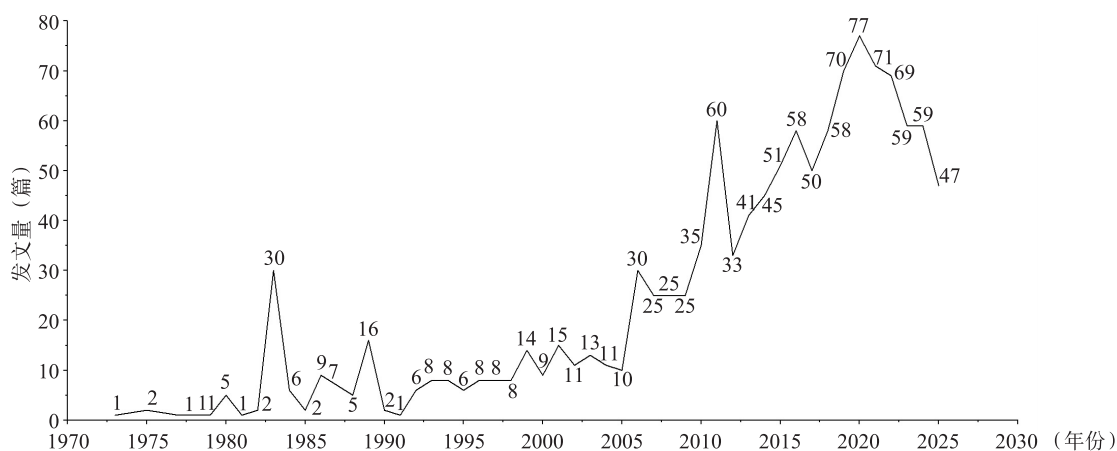


图1 引用文献数量折线图

引用文献的学科来源统计表明,管理学学科内引用文章数量高达615篇,占总文章数量的50.6%,表明管理学受到理论“有趣性”的影响最为深远。理论“有趣性”在管理学中的影响主要体现在教学培养、期刊评审与理论评价三个层面。首先,“有趣性”成为培养管理学学生的重要目标。Tihanyi (2020)指出Davis (1971)这篇文章成为大多数博士课程中的必读内容,此外Parboteeah (2012)在参与美国管理学会的博士学位论文评选比赛时,从Davis (1971)的“有趣性”视角肯定了获奖论文。同时,Swedberg (2014)、Podsakoff (2018)等为青年学者提供了相关的实践建议。其次,“有趣性”逐渐成为管理学期刊的重要评审标准。《美国管理学会评论》(Academy of Management Review, AMR)前任编辑Kilduff (2006)指出好的

理论论文应利用既有文献寻求“有趣性”,Barney (2018)亦提出“有趣性”是理论文章发表于管理学顶级期刊的可能原因;自20世纪90年代起,《美国管理学会评论》评审的主要原则之一是文章新颖性(Brief, 2003),《美国管理学会学报》编委调查显示57%的成员将“反直觉”视为一篇文章“最有趣”的原因之一(Bartunek et al., 2006)。最后,“有趣性”也成为评估管理学理论的重要特征。Corley & Gioia (2011)认为“令人惊讶的”是理论贡献的评价标准之一,Zahra和Newey (2009)也表达了相似的观点。

(三) 何必“有趣”: 理论“有趣性”的多方争议

理论“有趣性”所引起的各方争议,概括于表2。

表 2 理论“有趣性”的各方争议

争议阵营	代表学者	核心观点	主要关切
支持与认可	Sandberg; Alvesson; Chen 等 (2007); Gray 和 Cooper (2010); Bartunek 等 (2006); Brief (2003); Corley & Gioia (2011); Hensel (2021)	理论的价值在于挑战既有假设	确立反直觉作为理论贡献的主要标准
忧虑与警示	Pillutla 和 Thau (2013); Aksom 等 (2020)	过度追求有趣性会损害学术严谨性	防止理论建构脱离实践
质疑与批评	Ryan (2008); Tsang (2022); Wright (2024); Wright 等 (2024)	理论“有趣性”存在严重缺陷, 没有科学价值	主张面向现象与问题的科学研究
反思与平衡	Davis (1986); Corley 和 Gioia (2011); Salvador (2011); Tihanyi (2020)	应当平衡理论的有趣性与现实重要性	寻求理论“有趣性”与“重要性”之间的平衡点

资料来源：作者自制。

国外许多管理学学者支持与认可 Davis (1971) 关于理论“有趣性”的观点。一些学者支持和认可整体观点，代表性学者 Sandberg、Alvesson 多次讨论 Davis (1971) 的理论，认同研究者应挑战既有假设 (Sandberg & Alvesson, 2011; Alvesson & Sandberg, 2011, 2013, 2024)。而另一些学者则支持部分观点，如 Chen 等 (2007) 同意那些与受众基本假设完全一致的结果往往是无趣且显而易见的；Gray 和 Cooper (2010) 认可 Davis (1971) “看似是 X，实为非 X” 的公式。同时，Bartunek 等 (2006) 对期刊编委和媒体人士的调查为 Davis (1971) 的论点提供了实践支持，表明在实证研究中，挑战当前假设的文章极有可能是有趣的。此外，Brief (2003)、Corley 和 Gioia (2011) 等学者从期刊审稿的角度支持理论需要具有“有趣性”的特征。近年，Hensel (2021) 肯定了 Davis (1971) 对于所有学科的价值。因此，Davis (1971) 的研究收获了一众支持。

还有一些学者对理论“有趣性”的过度推

崇表示忧虑。Pillutla 和 Thau (2013) 指出，组织科学领域对理论“有趣性”的痴迷可能导致三个严重后果：研究不可复制、理论碎片化和脱离实践。他们警告，过度强调反直觉事实可能会导致有争议的理论发展和不可靠的数据收集。Aksom 等 (2020) 则认为如果学者们被迫要求发表既具有理论贡献，也符合有趣和可发表要求的论文时，通常我们既得不到理论进步，也得不到重要的实证研究成果。尽管这些学者指出了过度推崇理论“有趣性”可能会产生的严重后果，但他们并非对理论“有趣性”进行严肃批评和根本否定，相反，他们认可理论“有趣性”的价值和意义。

然而，以曾荣光 (Eric W. K. Tsang) 为代表的学者质疑和批评理论“有趣性”。应该看到，理论“有趣性”即便流传了 50 余年，从 1215 篇引用文章来看其学术影响也并不算极为广泛。这说明，它虽然长期被部分学者奉为经典，但其普遍性并未得到学界的广泛认同，甚至一些学者认为这是一部“邪教经典” (Ryan, 2008)。

2022年,曾荣光的一篇名为“*That’s Interesting! A Flawed Article Has Influenced Generations of Management Researchers*”的文章是对Davis (1971)的猛烈抨击。这篇文章直接将Davis (1971)的理论看作是有严重缺陷的理论,认为“有趣性不是一个好的科学理论的优点,因此其在科学中没有什么价值”(Tsang, 2022),鼓励研究者去解释现象与解决问题。这篇文章与此后Wright (2024)的研究都表明,崇尚理论“有趣性”不利于管理学作为一门社会科学学科持续发展。对此,Tsang特别指出打破理论“有趣性”的束缚将为更富有成效的研究开辟新的机会(Wright et al., 2024)。因此,理论“有趣性”受到了以曾荣光为代表的学者群体的质疑和批评。

更多的学者主张反思和平衡理论“有趣性”和理论“重要性”。1986年,作为理论“有趣性”的开创者,Davis (1986)发表了一篇题为“‘*That’s Classic!*’ The Phenomenology and Rhetoric of Successful Social Theories”的文章,指出要回应主要受众的现实关注点,理论“重要性”逐渐受到学者关注。进入21世纪后,Corley与Gioia (2011)认为理论除了具备“原创性”之外,还应具备“实用性”。典型代表学者Salvador (2011)提出研究人员应该关注那些在理论上和实践上都很有趣的问题,而不是过分关注预期结果的反直觉。2020年,《美国管理学会学报》主编Tihanyi (2020)强调“有趣性”和“重要性”之间要达成平衡,即需要研究那些不仅有趣而且与社会相关的问题,进一步推动了理论“有趣性”与“重要性”的兼容态势。因此,理论“有趣性”和

理论“重要性”之间的平衡是学者们始终在反思的问题。

(四) 何立“有趣”: 理论“有趣性”与情境

尽管存在诸多争议,理论“有趣性”在今天不仅没有过时,反而展现出持久的价值。首先,理论“有趣性”的核心思想是其并未过时的关键。Davis (1971)所倡导的“挑战受众既有假设”的思想,深刻揭示了知识生产机制,指出理论进步源于对特定认知基准的超越。其次,尽管文献引用有限,但理论“有趣性”对于所有学科都很有价值(Hensel, 2021)。这一文章对于学者如何撰写吸引受众阅读的文章很有帮助(Gartner, 2011)。最后,理论“有趣性”的相对性,正是其价值体现的关键。Davis (1971)在文末强调,一个理论的价值,必须置于其所挑战的特定时代下的常识来评判。理论“有趣性”挑战了当时人们所认为的理论“正确性”假设(Davis, 1971),其具有推动知识进步的价值。

学者们的批评主要集中于对理论“有趣性”的应用异化,而并未能撼动该理论自身的合法性。以曾荣光为代表的批评阵营其矛头本质上指向的是学术界对理论“有趣性”的痴迷与滥用,而并不是“否定既有假设”这一核心思想本身。事实上,曾荣光本人曾在论文中表明,自己在《组织研究》(*Organization Studies*)上的一项关于迷信与商业决策的研究(Tsang, 2004),之所以被重视和发表,其原因之一就是“有趣”(曾荣光, 2011)。这位曾严厉批判理论“有趣性”的学者,其得意之作恰恰因“有趣性”而获得发表与关

注。因此，这些争议更多的是对理论“有趣性”的应用矫正，反而证实了理论“有趣性”存在的合法性。

然而，Davis（1971）的理论“有趣性”局限于话语生产而忽视了社会建构性。假设的否定性形式，是通过话语的修饰生产理论“有趣性”，但是，理论的“有趣性”本质是一个社会集体建构的过程。它并非仅源于理论话语的形式创新，而是在特定共同体中，通过知识竞争、制度认可与话语互动被逐步确立的。曾被认定为具有理论“有趣性”的命题形式可能由于人工智能时代的知识共同体的变化而被视为无趣，而那些被认定为“无趣”的命题，通过人工智能时代研究视角和研究方法的创新，反而也可能展现出理论“有趣性”。在人工智能变革科研范式的今天，理论“有趣性”愈发依赖于人机协同的结果，能否对传统知识进行反思和重塑、能否通过新型科研制度的审查、能否在话语协商中获得集体认同成为重要标准。因此，人工智能时代理论“有趣性”的实现需要超越 Davis 的话语生产限定，探索理论“有趣性”的深刻内涵以及人工智能驱动的具体路径。

三、理论“有趣性”的启示：人工智能驱动的社会科学有趣性

（一）“人工智能驱动的社会科学研究”的范式扩展

人工智能是一种模拟人类的技术（Acemoglu & Restrepo, 2018；Chowdhury et al., 2022）。在1956年的达特茅斯会议（Dartmouth College Summer workshop）上人工智能首次被提出（陈凤仙，2022），主流技术发展经历了五个阶段，如图2所示。早期符号主义阶段（1950s—1970s）主要是基于符号逻辑规则进行知识表征（Post, 1943），此后专家系统阶段（1970s—1990s）通过知识库、推理机等模块实现逻辑推理（Newell & Simon, 1976）。机器学习阶段（1990s—2010s）以支持向量机、随机森林等机器学习算法为代表，实现从规则驱动到数据驱动的转变（Bishop, 2006）。深度学习阶段（2010s—2020s）以深度神经网络为核心，Transformer 架构的突破推动了大语言模型发展（Vaswani et al., 2017）。当前生成式人工智能阶段（2020s 至今）则基于生成对抗网络、扩散模型等技术实现创造性内容生成（Ho et al., 2020）。

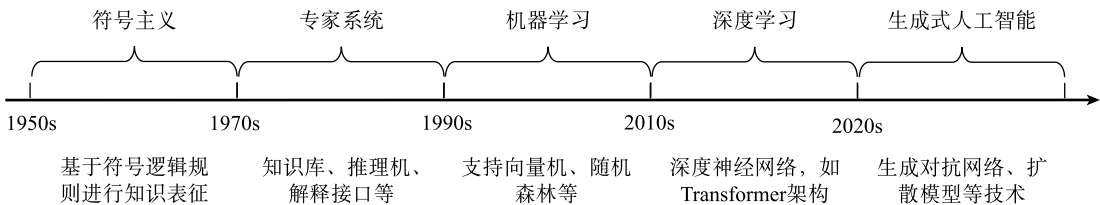


图2 人工智能技术发展阶段

科学研究经验范式、理论范式、计算范式、数据驱动范式，目前正迎来智能范式变革（王飞跃和缪青海，2023）。社会科学研究也

正在向人工智能驱动第五范式扩展，表3概括了具体的研究范式扩展。在研究方法上，人工智能正推动社会科学研究方法从理论驱动、

假设检验和因果推断，拓展至数据驱动、智能仿真和复杂因果分析（庞珣，2024；赵乃萱和李江，2025）。在研究对象上，人工智能自身作为技术成为社会科学的研究对象（Xu et al.，2024；马费成和陈帅朴，2024），同时人工智能不断影响社会行为、规范与价值体系，重塑人类社会秩序（李江，2024）。在研究范围与数据方面，研究从小样本走向大规模分析，并引入

语音、图像和视频等新型数据形态（政光景和吕鹏，2023；李江，2024）。同时，研究主体由单一人类研究者转向人机协同，研究组织呈现平台化趋势，学科边界不断交融（方晶等，2025；郑若婷等，2024）。在研究成果上人工智能不仅服务于论文与资政报告，也支持政策仿真、动态预测和决策支持等应用导向成果的生成（Zhu et al.，2026；祝忠明和寇蕾蕾，2025）。

表 3 “人工智能驱动的社会科学研究”的范式扩展

维度	旧范式	新范式
研究方法	理论驱动、假设检验、因果推断	数据驱动、智能仿真、复杂因果分析
研究对象	个体、组织、制度、文化	技术、行为、规范、价值
研究范围	小样本	全样本或大规模
研究数据	文本、问卷、实验、数据库	语音、图像、视频、传感器数据
研究主体	人类主体	人机协同
研究组织	分散化	平台化
研究边界	学科壁垒	跨学科融合
研究成果	学术论文、资政报告	政策仿真系统、动态预测模型、决策支持系统

资料来源：作者自制。

表 4 展示了“人工智能驱动的社会科学研究”中具有代表性的技术及其应用实践。大语言模型（Large Language Models, LLM）被用于文献综述、情感与话语分析及文本分类，Ji 等（2026）通过构建基于大语言模型的调查实验（LLM-based Survey Experiment）和实验室实验（LLM-based Laboratory Experiment）两个引文网络模型，展示了大语言模型在社会科学实验与模拟中的应用潜力。知识图谱（Knowledge Graphs, KG）主要服务于政策文本结构化与关系建模，Pan 等（2021）提出的智能调查系统体现了其在调查研究中的实践价值。文本挖掘（Text Mining, TM）以主题识别和文本量化为核

心，Blei 等（2003）的潜在狄利克雷分配模型（Latent Dirichlet Allocation, LDA）已成为社会科学话语分析的基础工具。基于主体的建模（Agent-based Modeling, ABM）用于政策仿真、市场模拟与社会演化研究，Jeon 等（2025）通过多主体模拟还原在线互动过程。深度因果学习（Deep Causal Learning, DCL）主要用于政策评估，Prasanna 等（2024）使用这一方法完成政策效果的识别。预测建模（Predictive Modeling, PM）则侧重于社会结果预测，Mullainathan 和 Spiess（2017）将其作为经济学中的预测工具。此外，多模态融合（Multimodal Fusion, MF）通过整合文本与图像等多源数据拓展了社会行为

分析（Decadri et al.，2025），而计算机视觉（Computer Vision，CV）则常被用于空间分析、
文史研究，例如利用卫星图像预测贫困水平（Jean et al.，2016）。

表 4“人工智能驱动的社会科学研究”的重点技术及应用实践

重点技术	常见应用	具体实践
大语言模型	文献综述、情感分析、话语分析、文本分类	Ji 等（2026）构建了两个基于大语言模型的社会科学引文网络模型
知识图谱	政策文本结构化、关系建模	Pan 等（2021）构建了基于知识图谱的社会科学智能调查系统
文本挖掘	主题识别、文本量化	Blei 等（2003）提出潜在狄利克雷分配模型，被广泛用于社会科学中的话语结构与议题分析
基于主体的建模	政策仿真、市场模拟、社会演化	Jeon et al.（2025）应用基于主体的建模来模拟社交媒体多主体在线对话互动
深度因果学习	政策评估、效应识别	Prasanna 等（2024）基于深度学习与反事实预测方法，评估政策的因果影响
预测建模	社会结果预测	Mullainathan 和 Spiess（2017）将预测建模用于计量经济学分析
多模态融合	社会行为分析	Decadri 等（2025）结合文本与图像数据分析社交媒体中的政治表达
计算机视觉	空间分析、文史研究	Jean 等（2016）基于卫星图像利用计算机视觉预测贫困水平

资料来源：作者自制。

（二）人工智能时代理论“有趣性”的社会建构性

理论“有趣性”是知识共同体在知识传统与制度情境下，通过话语协商达成的一种集体共识。Davis（1971）提出理论“有趣性”的总公式“看似是 X，实为非 X”，核心在于否定受众的既定假设。从认识论角度看，“否定”涉及否定方与被否定方，双方共同构成了知识共同体。“既定假设”是在知识传统与制度环境中被知识共同体建构并共享的认知。理论“有趣性”，并不是由否定方或被否定单方面决定，

也并不完全取决于知识的修辞包装，而是在知识传统与制度情境中由知识共同体开展话语协商而建构的。因此，“看似是 X，实为非 X”仅反映了话语层面的理论“有趣性”，而其本质是一种社会建构的结果。

作为“人工智能驱动的社会科学有趣性”的结果，“有趣的集体共识”在人机协同背景下呈现出特殊内涵与特征。其特殊内涵在于，它不再是人类知识共同体的单一共识，而是人类与人工智能通过持续话语协商达成的集体共识，既内含人类价值判断，也带有算法偏好。“有趣



的集体共识”具有三重特征。其一，生成主体的异质性，共识不再仅由人类学者产生，人工智能作为一个非人类主体者也参与到共识生产中（中国科学院创新范式研究组，2025）；其二，共识形成的加速性，人工智能表现出人类不可比拟的高效率，大幅压缩了“有趣的集体共识”的生产周期；其三，共识边界的开放性，人机协同突破了传统共同体的学科边界，使“有趣的集体共识”具有更强的跨学科适应力。因此，人工智能时代“有趣的集体共识”已因人机协同而发生了深刻的转变。

理论“有趣性”的社会建构依赖于知识传统、制度情境以及话语协商，人工智能正在变革这些方面。其一，理论“有趣性”的建构基础是知识传统。人工智能正在加速知识传统习得与挖掘，典型例子是大语言模型在文献方面表现出高效的整合和分析能力（Ji et al., 2026）。其二，理论“有趣性”须在制度情境中获得合法性认可。人工智能能够发挥模拟制度情境的作用，典型例子是 ChatGPT（Chat Generative Pre-trained Transformer）在同行评审模拟中的应用，Liang 等（2024）研究发现，57.4% 的受访研究者认为 GPT-4 生成的盲评意见总体有帮助。其三，理论“有趣性”在话语协商中生成。人工智能正在作为新的话语主体增强理论“有趣性”，典型例子是人工智能促进有趣的研究想法的生成，Si 等（2025）通过盲评人类研究者与人工智能生成的研究想法，发现后者的新颖性显著优于前者。

理论“有趣性”的核心过程是知识共同体“否定方”与“被否定方”的话语协商。一方面，知识共同体早期通过与知识传统、制度情

境进行互动，形成内部成员共享的假设。另一方面，共同体内部的“否定方”通过话语生产挑战这些既定假设，制造“看似是 X，实为非 X”的认知张力（Davis, 1971），“被否定方”则依据既有认知对挑战作出话语反馈。若话语反馈为“显而易见”“无关紧要”或“荒谬”（Davis, 1971），则理论归于“无趣”。若话语反馈为“真是有趣”（Davis, 1971），则知识共同体达成理论“有趣性”的集体建构与认同。因此，理论“有趣性”是知识共同体在历史与制度约束下，通过集体协商达成的一种共识。人工智能时代的理论“有趣性”，本质上仍然是社会集体建构和认同的结果。

（三）人工智能驱动的社会科学有趣性

基于 Davis（1971）及许多学者对理论“有趣性”实现方式的讨论，理论“有趣性”主要来源于三种假设。一是本质性假设，Davis（1971）提出的组织、构成、抽象等方式，Davis（1999）的概念魅力、Alvesson 和 Kärreman（2007）的神秘隐喻等都是通过挑战人们对现象基本成分、范畴或存在方式的固有认知实现理论“有趣性”。二是价值性假设，Davis（1971）的功能与评价方式、Sparrowe 和 Mayer（2011）的对比方法、Gray 和 Wegner（2013）的“令人惊讶”都是旨在通过颠覆关于现象功能或效用的既定评价来生成理论“有趣性”。三是关系性假设，Davis 的关联、协变等方式，Salvato 和 Aldrich（2012）的关注变量、Božič（2020）的多源数据与异质视角等方法都聚焦于通过挑战现象间的关联认知来产生理论“有趣性”。具体如表 5 所示。

表 5 理论“有趣性”的三种假设方式

假设类别	核心逻辑	Davis (1971) 的有趣方式	其他学者的有趣方式
本质性假设	人们对现象成分、范畴或存在方式的固有认知	组织、构成、抽象、泛化、稳定、共存	概念魅力 (Davis, 1999)、神秘隐喻 (Alvesson & Kärreman, 2007)
价值性假设	人们对现象功能或效用的固有认知	功能、评价	对比 (Sparrowe & Mayer, 2011)、“令人惊讶” (Gray & Wegner, 2013)
关系性假设	人们对现象之间相互联系或作用方式的固有认知	关联、协变、对立、因果	关注变量 (Salvato & Aldrich, 2012)、多源数据与异质视角 (Božič, 2020)

资料来源：作者自制。

表 6 展示了人工智能驱动社会科学理论“有趣性”的具体实现路径，即人类与人工智能组成的知识共同体通过“扎根知识传统—适应制度情境—话语协商有趣”建构“有趣的集体共识”。其中，社科学者发挥关注假设、重要化

理论、引导与叙事的作用，而人工智能则发挥加速挖掘、模拟预测、处理实施与叙事优化的作用，二者组成的知识共同体推动建构“有趣的集体共识”。

表 6 人工智能驱动社会科学理论“有趣性”的具体路径

知识共同体		主要相关的人工智能技术	
	人类（主要指社科学者）	人工智能	
第一步 扎根知识传统			
本质性假设	关注被视为理所当然、不言自明的假设 (Davis, 1971)	加速挖掘	大语言模型、文本挖掘
价值性假设	关注评价高度一致的假设 (Davis, 1971)		文本挖掘、预测建模
关系性假设	关注认知高度一致的假设 (Davis, 1971)		知识图谱、深度因果学习
第二步 适应制度情境			
评审制度、资助制度等	做重要的理论	模拟预测	大语言模型、基于主体的建模、预测建模
第三步 话语协商有趣			
有趣的语料	引导	处理	多模态融合、计算机视觉
有趣的方法		实施	基于主体的建模、预测建模
否定—被否定	否定方 被否定方		大语言模型、基于主体的建模
有趣的言语	“真是有趣!” 的理论叙事	叙事优化	大语言模型、计算机视觉
结果：有趣的集体共识			

资料来源：作者自制。

第一步，扎根知识传统。社科学者需关注 Davis (1971) 提出的三类根本性假设：被视为理所当然的“本质性假设”、评价高度一致的“价值性假设”，以及认知高度一致的“关系性假设”。相应地，人工智能在此步骤中主要发挥加速挖掘的作用，主要涉及以下相关的人工智



能技术。利用“大语言模型”和“文本挖掘”可以分析教科书、论文等文本中不言自明的预设；通过“文本挖掘”和“预测建模”可以处理涉及价值判断的一致性问题的；借助“知识图谱”和“深度因果学习”可以揭示和验证变量间的常见但可能产生反转的关系。由此，通过扎根知识传统为理论奠定坚实的知识基础。

第二步，适应制度情境。要适应具体的制度情境，社科学者就必须致力于做重要的理论。在评审制度、资助制度等多重约束下，单纯追求理论“有趣性”往往不足以适应制度的筛选标准，《美国管理学会学报》主编 Tihanyi (2020) 呼吁学者在追求“有趣”的同时，更要关注“重要”。理论“重要性”体现为其对社会实践的回应程度，社科学者要创造与现实世界密切相关的理论 (Geddes, 2003)，以研究是否“天然重要”为重要动机 (陈硕, 2024)。人工智能在此阶段主要发挥模拟预测功能，为理论是否契合制度情境提供判断工具。主要技术有“大语言模型”“基于主体的建模”和“预测建模”，能够模拟制度运行、预测理论的重要影响和接受度，从而帮助学者增强理论“重要性”，使其更符合制度要求与期待。由此，通过增强理论“重要性”适应制度情境。

第三步，话语协商有趣。理论“有趣性”最终需要在话语协商中获得集体认同。首先，社科学者引导人工智能扩展语料范围，使理论不再仅依赖经典文本，而是可以收集多模态证据形成“有趣的语料”，如可采用“多模态融合”和“计算机视觉”技术来处理文本、图像等多样态数据。其次，“有趣的方法”可以依赖人工智能来实施，如新兴方法“基于主体的建

模”“预测建模”等，社科学者在其中发挥引导作用。接着，话语协商在“否定方”与“被否定方”之间的“否定—被否定”互动中展开。其中，人工智能相关技术，“大语言模型”可辅助生成对立论点、“基于主体的建模”可模拟多元学术对话，由此人类与人工智能在不断的否定与被否定中建构理论。最后，社科学者主导能引发“真是有趣！”共鸣的理论叙事，人工智能则承担“叙事优化”的工作，通过“大语言模型”和“计算机视觉”来润色和呈现故事，最终人机协同完成“有趣的言语”构建。由此，不同主体建构和达成“有趣”的理论共识。

第四步，整个路径导向“有趣的集体共识”。这意味着，通过社科学者与人工智能的协同合作，从深度扎根知识传统中的理论假设，到灵活适配评审、资助等现实制度情境，再到反复进行以“否定—被否定”为核心的话语协商，理论“有趣性”得以在知识共同体中集体建构，从而形成一种共享的理解。由此，理论“有趣性”是在人类与人工智能协同下，被不断建构与认同的集体共识。

(四) 在实践中应当警惕的风险

第一，在“扎根知识传统”阶段，需警惕“人工智能幻觉” (AI Hallucinations)。“人工智能幻觉”指模型基于概率生成机制，输出形式上流畅、具有高度确定性，实则违背事实或逻辑的内容的表象 (OpenAI, 2023)。在社会科学的知识传统中，理论“有趣性”的构建依赖于真实发生的“事实”，而人工智能可以像人类一样产生故意欺骗行为 (Hutson, 2024)，这直接威胁到理论“有趣性”所扎根的知识的可靠性。“人工智能幻觉”可能生成错误知识，通过巧妙

缝合虚假关联或曲解语境，制造出看似深刻、实则歪曲“事实”的知识，误导社科学者的判断。因此，社科学者必须建立严格的验证机制，不能将人工智能输出直接等同于理论“有趣性”的有效素材，而需对其内容进行多次审视和事实核查，保证理论“有趣性”扎根于可靠的知识传统。

第二，在“适应制度情境”阶段，需警惕人工智能造成的“算法趋同”。人工智能正在深刻改变评审、资助等外部规则体系（Lin, 2024），这可能使理论“有趣性”的评价标准被技术隐形固化。当人工智能普遍介入论文评审，其内置的算法偏好可能强化某种理论“有趣性”模板，导致学术产出趋同，削弱了由学者个人兴趣（Renninger, 2000）、早期学术印记（Marquis & Tilesik, 2013）、不同研究共同体（Salvato & Aldrich, 2012）以及不同受众背景（Bartunek et al., 2006）所孕育的“有趣”多样性。最终，人工智能的渗透，可能使理论“有趣性”的集体协商沦为对算法标准的适应性策略。因此，必须警惕评审、资助等制度过度依赖人工智能，设计以人类为主体的评价机制，以保护多元理论的趣味性。

第三，在“话语协商有趣”阶段，需警惕人工智能带来的“批判惰性”。人工智能正在改变话语协商的材料、方式与成果，但也可能导致协商过程流于表面。对现象进行判断（Alvesson & Nørmark, 2026）、质疑（Alvesson & Sandberg, 2024），是社科学者批判性思维的体现。若学者过度依赖人工智能生成的流畅叙事，其主动质疑、深度思辨的能力面临退化风险，有可能使得话语协商异化为对人工智能输出内容

的简单认可或修饰。因此，在与人工智能进行话语协商时，社会科学学者必须保持自身在话语协商中的主导地位，将人工智能作为辅助工具，持续运用批判性思维去主动挑战既定假设，推动理论“有趣性”生成。

四、结论与展望

面向国内社会科学研究需要增强有趣性的现状，本文对理论“有趣性”进行了系统综述，并提出人工智能驱动社会科学理论“有趣性”的实现路径。文章首先综述 Davis（1971）提出的理论“有趣性”，梳理了该理论的主要内容、发展脉络与学科影响以及多方争议，并进行了研究述评；其次回顾“人工智能驱动的社会科学研究”的范式扩展，指出人工智能时代理论“有趣性”的社会建构性；最后提出社科学者应当如何借助人工智能驱动社会科学理论“有趣性”，提示应当警惕的三大风险。

本文的贡献包括以下几点。其一，为增强当前国内社会科学理论“有趣性”提供人工智能驱动的方法论路径，引导社科学者主动运用人工智能通过“扎根知识传统—适应制度情境—话语协商有趣”建构“有趣的集体共识”；其二，突破理论“有趣性”窄化于话语生产的局限，揭示理论“有趣性”的社会建构性，指出理论“有趣性”是知识共同体在知识传统与制度情境下，通过话语协商达成的一种集体共识；其三，试图弥合经典社会科学哲学与当代人工智能方法论之间的理论断层，从而推动社会科学研究在理论与实践上双向深化。

未来研究需要深入探究人工智能驱动社会



科学理论“有趣性”的审查制度。随着人工智能驱动理论“有趣性”的能力不断增强,其自动化、规模化和加速研究的特性,也可能导致大量缺乏实证基础、逻辑薄弱甚至具有误导性的“伪有趣”内容涌入学术领域。因此,建立有效审查制度至关重要。同时应该明确,审查制度的目标并不是抑制有趣,而是引导社会科学学者借助人工智能服务于既具有科学趣味又具有现实价值的知识生产。

接受编辑:贾良定

收稿时间:2025年10月7日

接收时间:2026年5月21日

作者简介

周贤永,西南交通大学公共管理学院副院长、副教授、硕士研究生导师,于西南交通大学获得管理学博士学位。研究兴趣包括公共管理、科技创新管理和数字治理。研究成果发表在《科学学与科学技术管理》《软科学》等期刊上。

刘玉欣(通讯作者, E-mail: liuyuxin254@mails.ucas.ac.cn),中国科学院大学公共政策与管理学院硕士研究生。研究兴趣包括人工智能、数智治理和公共服务。

参考文献

- [1] 白长虹、林雪娇、刘晓:《面向人机共生的未来:人工智能驱动管理学研究的逻辑与证据》,《南开管理评论》,2025年第12期。
- [2] 白胜:《有趣性:管理研究的新策略》,《外国经济与管理》,2014年第5期。
- [3] 白胜:《聚焦反常:选取管理研究问题的新策

略》,《科技进步与对策》,2015年第17期。

[4] 白胜、严茜:《有趣性会计教学过程及其实施》,《商业会计》,2015年第16期。

[5] 陈凤仙:《人工智能发展水平测度方法研究进展》,《经济学动态》,2022年第2期。

[6] 陈硕:《非此即彼:社会科学研究指南》,团结出版社2024年版。

[7] 陈晓萍、沈伟(主编):《组织与管理研究的实证方法(第三版)》,北京大学出版社2018年版。

[8] 方晶、陈新、黄帅:《人工智能技术与哲学社会科学组织方式创新——以哲学社会科学实验室为例》,《浙江社会科学》,2025年第8期。

[9] 贾良定、林泽民、王涛、卜濛濛、何刚:《P值之争与管理学研究:先验概率的意义》,《管理科学学报》,2024年第3期。

[10] 李江:《AI影响了你的社会科学研究吗?》,《图书情报知识》,2024年第5期。

[11] 李智超、刘鹏、韩啸、甘甜:《生成式人工智能与社会科学方法论》,《公共管理与政策评论》,2026年第1期。

[12] 马费成、陈帅朴:《生成式人工智能赋能哲学社会科学研究》,《武汉大学学报(哲学社会科学版)》,2024年第6期。

[13] 马亮:《人工智能是万灵药吗?——人机协同的社会科学研究何以可能》,《天府新论》,2025年第6期。

[14] 马亮:《学术祛魅:实证研究十讲》,中国人民大学出版社2025年版。

[15] 庞珣:《人工智能赋能社会科学研究探析——生成式行动者、复杂因果分析与人机科研协同》,《世界经济与政治》,2024年第7期。

[16] 盛昭瀚:《问题导向:管理理论发展的推动力》,《管理科学学报》,2019年第5期。

[17] 王飞跃、缪青海:《人工智能驱动的科学新范式:从AI4S到智能科学》,《中国科学院院刊》,

2023 年第 4 期。

[18] 谢耘耕、张伶俐：《AIGC 技术驱动下舆论学研究方法的范式重塑与路径探索》，《新媒体与社会》，2025 年第 1 期。

[19] 曾荣光：《十字路口的中国管理研究：一些哲学层面的思考》，《重庆大学学报（社会科学版）》，2011 年第 4 期。

[20] 张志学、鞠冬、马力：《组织行为学研究的现状：意义与建议》，《心理学报》，2014 年第 2 期。

[21] 赵乃萱、李江：《机遇与挑战：人工智能时代人文社会科学知识生产范式变革》，《自然辩证法研究》，2025 年第 12 期。

[22] 政光景、吕鹏：《生成式人工智能与哲学社会科学新范式的涌现》，《江海学刊》，2023 年第 4 期。

[23] 郑若婷、于文轩、赵昊雪、肖依娜：《“AI 驱动的社会科学研究与公共治理新范式的构建” 高端学术论坛综述》，《公共管理学报》，2024 年第 1 期。

[24] 中国科学院创新范式研究组：《加快人工智能驱动的知识创新体系建设》，《中国科学院院刊》，2025 年第 10 期。

[25] 祝忠明、寇蕾蕾：《AI 驱动人文社会科学范式升维：从工具依赖到知识共生》，《图书与情报》，2025 年第 3 期。

[26] Acemoglu, D. , & Restrepo, P. 2018. The race between man and machine: Implications of technology for growth, factor shares, and employment. *American Economic Review*, 108: 1488 – 1542.

[27] Aksom, H. , Zhylynska, O. , & Gaidai, T. 2020. Can institutional theory be refuted, replaced or modified? *International Journal of Organizational Analysis*, 28: 135 – 159.

[28] Alvesson, M. , & Kärreman, D. 2007. Constructing mystery: Empirical matters in theory development. *Academy of Management Review*, 32: 1265 – 1281.

[29] Alvesson, M. , & Nørmark, D. 2026. *Return to*

judgement: The case for post – infantilized management. Bristol: Bristol University Press.

[30] Alvesson, M. , & Sandberg, J. 2011. Generating research questions through problematization. *Academy of Management Review*, 36: 247 – 271.

[31] Alvesson, M. , & Sandberg, J. 2013. *Constructing research questions: Doing interesting research*. Thousand Oaks, CA: Sage.

[32] Alvesson, M. , & Sandberg, J. 2024. The art of phenomena construction: A framework for coming up with research phenomena beyond “the usual suspects”. *Journal of Management Studies*, 61: 1737 – 1765.

[33] Barney, J. B. 2018. Editor’s comments: Positioning a theory paper for publication. *Academy of Management Review*, 43: 345 – 348.

[34] Bartunek, J. M. , Rynes, S. L. , & Ireland, R. D. 2006. What makes management research interesting and why does it matter? *Academy of Management Journal*, 49: 9 – 15.

[35] Bishop, C. M. 2006. *Pattern recognition and machine learning*. New York: Springer.

[36] Blei, D. M. , Ng, A. Y. , & Jordan, M. I. 2003. Latent Dirichlet allocation. *Journal of Machine Learning Research*, 3: 993 – 1022.

[37] Božić, B. , Siebert, S. , & Martin, G. 2020. A grounded theory study of factors and conditions associated with customer trust recovery in a retailer. *Journal of Business Research*, 109: 440 – 448.

[38] Brekhus, W. 1998. A sociology of the unmarked: Redirecting our focus. *Sociological Theory*, 16: 34 – 51.

[39] Brief, A. 2003. Editor’s comments: AMR—The often misunderstood journal. *Academy of Management Review*, 28: 7 – 8.

[40] Carr, T. H. 1983. Event structure, interest,

importance, and coherence: Where does point theory fit? *Behavioral and Brain Sciences*, 6: 597 – 598.

[41] Chen, X. – P., Wasti, S. A., & Triandis, H. C. 2007. When does group norm or group identity predict cooperation in a public goods dilemma? The moderating effects of idiocentrism and allocentrism. *International Journal of Intercultural Relations*, 31: 259 – 276.

[42] Chowdhury, S., Budhwar, P., Dey, P. K., Joel – Edgar, S., & Abadie, A. 2022. AI – employee collaboration and business performance: Integrating knowledge – based view, socio – technical systems and organisational socialisation framework. *Journal of Business Research*, 144: 31 – 49.

[43] Corley, K. G., & Gioia, D. A. 2011. Building theory about theory building: What constitutes a theoretical contribution? *Academy of Management Review*, 36: 12 – 32.

[44] Davis, M. S. 1971. That's interesting!: Towards a phenomenology of sociology and a sociology of phenomenology. *Philosophy of the Social Sciences*, 1: 309 – 344.

[45] Davis, M. S. 1986. “That's classic!” The phenomenology and rhetoric of successful social theories. *Philosophy of the Social Sciences*, 16: 285 – 301.

[46] Davis, M. S. 1999. Aphorisms and clichés: The generation and dissipation of conceptual charisma. *Annual Review of Sociology*, 25: 245 – 269.

[47] Decadri, S., Mancosu, M., Negri, F., & Scaduto, G. 2025. Aligning emotions: A comparative analysis of text and imagery by European party leaders on Instagram. *The International Journal of Press/Politics*, Online First. <https://doi.org/10.1177/19401612251339107>.

[48] Dutton, J. E., & Dukerich, J. M. 2006. The relational foundation of research: An underappreciated dimension of interesting research. *Academy of Management Journal*, 49: 21 – 26.

[49] Gartner, W. B. 2011. When words fail: An entrepreneurship glossolalia. *Entrepreneurship & Regional Development*, 23: 9 – 21.

[50] Geddes, B. 2003. *Paradigms and sand castles: Theory building and research design in comparative politics*. Ann Arbor, MI: University of Michigan Press.

[51] Gray, K., & Wegner, D. M. 2013. Six guidelines for interesting research. *Perspectives on Psychological Science*, 8: 549 – 553.

[52] Gray, P. H., & Cooper, W. H. 2010. Pursuing failure. *Organizational Research Methods*, 13: 620 – 643.

[53] Hafermalz, E., Johnston, R. B., Hovorka, D. S., & Riemer, K. 2020. Beyond “mobility”: A new understanding of moving with technology. *Information Systems Journal*, 30: 762 – 786.

[54] Hao, Q., Xu, F., Li, Y., & Evans, J. 2026. Artificial intelligence tools expand scientists' impact but contract science's focus. *Nature*, 649: 1237 – 1243.

[55] Hensel, P. G. 2021. Reproducibility and replicability crisis: How management compares to psychology and economics—A systematic review of literature. *European Management Journal*, 39: 577 – 594.

[56] Ho, J., Jain, A., & Abbeel, P. 2020. Denoising diffusion probabilistic models. *Advances in Neural Information Processing Systems*, 33: 6840 – 6851.

[57] Hutson, M. 2024. Two – faced AI language models learn to hide deception. *Nature*, January 23. Retrieved from <https://doi.org/10.1038/d41586-024-00189-3> (accessed September 5, 2026).

[58] Jean, N., Burke, M., Xie, M., Davis, W. M., Lobell, D. B., & Ermon, S. 2016. Combining satellite imagery and machine learning to predict poverty. *Science*, 353: 790 – 794.

[59] Jeon, M. S., Mendoza, M., Fernández, M.,

- Providel, E. , Rodríguez, F. , Espina, N. , Carvallo, A. , & Abeliuk, A. 2025. Simulating conversations on social media with generative agent – based models. *EPJ Data Science*, 14: 79.
- [60] Ji, J. , Lei, R. , Pan, X. , Wei, Z. , Sun, H. , Lin, Y. , Chen, X. , Yang, Y. , Li, Y. , Ding, B. , & Wen, J. – R. 2026. Leveraging LLM – based agents for social science research: Insights from citation network simulations. *Humanities and Social Sciences Communications*, 13: 127.
- [61] Kilduff, M. 2006. Editor’s comments: Publishing theory. *Academy of Management Review*, 31: 252 – 255.
- [62] Kirchherr, J. 2023. Bullshit in the sustainability and transitions literature: A provocation. *Circular Economy and Sustainability*, 3: 167 – 172.
- [63] Leatherbee, M. , & Katila, R. 2020. The lean startup method: Early – stage teams and hypothesis – based probing of business ideas. *Strategic Entrepreneurship Journal*, 14: 570 – 593.
- [64] Liang, W. , Zhang, Y. , Cao, H. , Wang, B. , Ding, D. Y. , Yang, X. , Vodrahalli, K. , He, S. , Smith, D. S. , Yin, Y. , McFarland, D. A. , & Zou, J. 2024. Can large language models provide useful feedback on research papers? A large – scale empirical analysis. *NEJM AI*, 1.
- [65] Lin, Z. 2024. Towards an AI policy framework in scholarly publishing. *Trends in Cognitive Sciences*, 28: 85 – 88.
- [66] Locke, K. , & Golden – Biddle, K. 1997. Constructing opportunities for contribution: Structuring intertextual coherence and “problematizing” in organizational studies. *Academy of Management Journal*, 40: 1023 – 1062.
- [67] Marquis, C. , & Tilcsik, A. 2013. Imprinting: Toward a multilevel theory. *Academy of Management Annals*, 7: 195 – 245.
- [68] Mullainathan, S. , & Spiess, J. 2017. Machine learning: An applied econometric approach. *Journal of Economic Perspectives*, 31: 87 – 106.
- [69] Newell, A. , & Simon, H. A. 1976. Computer science as empirical inquiry: Symbols and search. *Communications of the ACM*, 19: 113 – 126.
- [70] OpenAI. 2023. *GPT – 4 technical report*. Retrieved from [https://cdn.openai.com/papers/gpt – 4.pdf](https://cdn.openai.com/papers/gpt-4.pdf) (accessed May 23, 2026) .
- [71] Pan, J. Z. , Edelstein, E. , Bansky, P. , & Wyner, A. 2021. A knowledge graph based approach to social science surveys. *Data Intelligence*, 3: 477 – 506.
- [72] Parboteeah, K. P. 2012. Social issues in management division dissertation award competition for 2011. *Business & Society*, 51: 650 – 658.
- [73] Pillutla, M. M. , & Thau, S. 2013. Organizational sciences’ obsession with “that’s interesting!”: Consequences and an alternative. *Organizational Psychology Review*, 3: 187 – 194.
- [74] Podsakoff, P. M. , Podsakoff, N. P. , Mishra, P. , & Escue, C. 2018. Can early – career scholars conduct impactful research? Playing “small ball” versus “swinging for the fences” . *Academy of Management Learning & Education*, 17: 496 – 531.
- [75] Prasanna, A. N. , Grecov, P. , Weng, A. D. , & Bergmeir, C. 2024. Causal effect estimation with global probabilistic forecasting: A case study of the impact of COVID – 19 lockdowns on energy demand. *IEEE Transactions on Power Systems*, 39: 3417 – 3430.
- [76] Renninger, K. A. 2000. Individual interest and its implications for understanding intrinsic motivation. In C. Sansone & J. M. Harackiewicz (Eds.), *Intrinsic and extrinsic motivation: The search for optimal motivation and performance*. CA: Academic Press.
- [77] Ryan, D. 2008. Murray S. Davis 1940 – 2007. *Footnotes*, 36: 14 – 15.

- [78] Salvador, F. 2011. On the importance of good questions and empirically grounded theorizing. *Journal of Supply Chain Management*, 47: 21 – 22.
- [79] Salvato, C. , & Aldrich, H. E. 2012. “That’s interesting!” in family business research. *Family Business Review*, 25: 125 – 135.
- [80] Sandberg, J. , & Alvesson, M. 2011. Ways of constructing research questions: Gap – spotting or problematization? *Organization*, 18: 23 – 44.
- [81] Shortliffe, E. H. 1976. *Computer – based medical consultations: MYCIN*. New York: Elsevier.
- [82] Si, C. , Yang, D. , & Hashimoto, T. 2025. Can LLMs generate novel research ideas? A large – scale human study with 100 + NLP researchers. In *The Thirteenth International Conference on Learning Representations*.
- [83] Sparrowe, R. T. , & Mayer, K. J. 2011. Publishing in AMJ—Part 4: Grounding hypotheses. *Academy of Management Journal*, 54: 1098 – 1102.
- [84] Swedberg, R. 2014. *The art of social theory*. Princeton, NJ: Princeton University Press.
- [85] Tihanyi, L. 2020. From “That’s interesting” to “That’s important” . *Academy of Management Journal*, 63: 329 – 331.
- [86] Trigg, R. 1985. *Understanding social science: A philosophical introduction to the social sciences*. Oxford: Basil Blackwell.
- [87] Tsang, E. W. K. 2004. Toward a scientific inquiry into superstitious business decision – making. *Organization Studies*, 25: 923 – 946.
- [88] Tsang, E. W. K. 2022. That’s interesting! A flawed article has influenced generations of management researchers. *Journal of Management Inquiry*, 31: 150 – 164.
- [89] Van Quaquebeke, N. , Tonidandel, S. , & Banks, G. C. 2025. Beyond efficiency: How artificial intelligence (AI) will reshape scientific inquiry and the publication process. *The Leadership Quarterly*, 36: 101895.
- [90] Vaswani, A. , Shazeer, N. , Parmar, N. , Uszkoreit, J. , Jones, L. , Gomez, A. N. , Kaiser, Ł. , & Polosukhin, I. 2017. Attention is all you need. *Advances in Neural Information Processing Systems*, 30: 5998 – 6008.
- [91] Wright, T. A. 2024. Truth and the science of management: Does the “emperor” still have any clothes? *Journal of Management Inquiry*, 33: 3 – 10.
- [92] Wright, T. A. , Emich, K. , Pearce, J. L. , Ramoglou, S. , Ashkanasy, N. M. , Bartunek, J. M. , Kunisch, S. , Denyer, D. , Foss, N. J. , Klein, P. G. , Town, S. , Hollwitz, J. , Barney, C. E. , Harms, P. , Munyon, T. P. , Seijts, G. , & Tsang, E. W. K. 2024. Advocacy and the search for truth in management scholarship: Can the Twain ever meet? *Journal of Management Inquiry*, 33: 11 – 25.
- [93] Xu, R. , Sun, Y. , Ren, M. , Guo, S. , Pan, R. , Lin, H. , Sun, L. , & Han, X. 2024. AI for social science and social science of AI: A survey. *Information Processing & Management*, 61: 103665.
- [94] Zahra, S. A. , & Newey, L. R. 2009. Maximizing the impact of organization science: Theory – building at the intersection of disciplines and/or fields. *Journal of Management Studies*, 46: 1059 – 1075.
- [95] Zhu, Y. , Lu, Y. , Xie, H. , Ye, J. , & Chen, M. 2026. A quasi – experimental analysis of capabilities and limitations of generative AI in academic content evaluation in social sciences. *Information Processing & Management*, 63: 104365.