

• 研究构想(Conceptual Framework) •

领导对员工—生成式人工智能建议的反应机制： 基于社会比较视角多层次研究*

韩 翼¹ 马朝翊¹ 宗树伟²

(¹中南财经政法大学工商管理学院, 武汉 430073) (²西南石油大学经济管理学院, 成都 610500)

摘 要 GenAI 参与组织决策已经成为不可阻挡的趋势, 但学术界对于 GenAI 建议采纳的研究尚不完善。在领导决策过程中, 领导如何比较员工建议、GenAI 建议、员工—GenAI 团队建议并进行采纳? 领导的感知有何差异? 为此, 本研究从社会比较理论的视角出发, 通过 5 个子课题的研究, 主要解决以下 5 个问题: 1) GenAI 建议采纳如何进行界定? 2) 领导对于员工—GenAI 建议采纳存在何种差异? 3) 领导对于员工—GenAI 团队建议采纳的差异如何? 4) 如何比较员工—GenAI 建议采纳的差异效果? 5) 员工—GenAI 协同过程中建议障碍如何差异化影响建议质量? 干预策略是否有效? 最终, 本研究系统构建了领导对于员工—GenAI 建议采纳的多层次模型, 为跨学科理论发展赋予了新的视角, 也为组织优化人智协同决策, 降低技术风险提供了实践指导。

关键词 领导力, GenAI 建议, 建议反应, 社会比较理论, 感知差异

分类号 B849: C93

1 问题提出

以 ChatGPT 为代表的生成式人工智能(GenAI)的迅猛发展正对组织产生日益显著的影响(刘伟, 谭文辉, 2025; Rizzo et al., 2024)。在工作场所中, GenAI 正被引入那些曾完全依赖于人类判断的决策情境中(Kahr et al., 2024), 通过基于大型数据集和机器学习提供的建议, 正在深刻影响着领导决策, 并助力管理决策的优化(Mahmud et al., 2023; 许丽颖 等, 2025)。

GenAI 在多个组织领域开辟了新的可能性, 包括临床护理(Jeblick et al., 2024)、教育(Kasneci et al., 2023)、艺术和音乐(Civit et al., 2022; Oksanen et al., 2023)以及设计(Jiang et al., 2023)。与依赖历史数据进行预测的 AI 不同, GenAI 基于海量预训练数据集生成新的内容(Feuerriegel et al.,

2024), 并结合人类反馈来优化其算法。这一特性使其能够跨文本、图像、音频、编程代码、模拟及视频等多种模态合成与情景相关的新颖内容(Tilton et al., 2023), 而此类生成的内容、推荐或建议能够对领导决策产生积极影响(Huang et al., 2024)。

要将 GenAI 技术引入建议领域, 以推动领导采纳其建议或“员工—GenAI”混合建议, 必须厘清其背后的促进或抑制机制。具体而言, 领导如何评估或比较员工建议、GenAI 建议及员工—GenAI 团队建议, 进而做出采纳决策, 是一个具有重要研究价值的问题, 具体可分述如下:

首先, 如何界定 GenAI 建议采纳? 目前学术界对此仍存分歧。譬如有的学者将其简单定义为个体使用了 GenAI 的相关产品(Agrawal et al., 2024), 另外一些学者则认为 GenAI 建议采纳是一种评价后选择接受的决策过程(宗树伟 等, 2025)。尽管这些不同的定义各有其价值, 但概念的差异不仅不利于深入挖掘 GenAI 建议采纳的重要内涵, 还可能阻碍后续研究的推进。鉴于此, 本研究聚焦于领导在员工—GenAI 建议决策中的核

收稿日期: 2025-09-09

* 国家自然科学基金面上项目(72572169); 国家自然科学基金青年项目(72402183)。

通信作者: 韩翼, E-mail: hanyi7009@163.com

心情境,结合社会比较理论,将 GenAI 建议采纳界定为:领导在接收 GenAI (或员工-GenAI 协同团队)建议后,通过与员工建议进行多维度比较,形成认知层面的感知和情感反应,并最终将建议内容整合到组织决策方案中或直接实施其核心措施的一种动态过程。这一定义主要强调决策者“先评价,后选择”的动态过程,也契合领导在人智协同中权衡、反应乃至实施的现实决策逻辑。

第二,领导对于员工建议与 GenAI 建议的采纳存在何种差异?学术界对此呈现出两极分化的观点。一些学者认为,由于 GenAI 建议存在可解释性不足、准确性争议及领导的信任缺失等问题,这可能引发诸多阻碍,导致领导更倾向于抵制 GenAI 建议,转而信任员工建议(Mahmud et al., 2023)。另一些学者则认为,即使人智建议都不完美,二者合成后仍能提升决策效果(Sachin & Schecter, 2024; Choudhary et al., 2025)。此外,在工作场所中采纳 GenAI 建议有助于减少机械性、重复性工作,从而助力领导提升决策效能(You et al., 2022)。值得注意的是,组织中 GenAI 建议采纳的影响机制可能区别于传统领导纳谏,这一差异取决于领导对于员工建议与 GenAI 建议的差异感知,具体包括领导感知到的建议来源或目标的比较、维度比较和动机比较(Matthews & Kelemen, 2025)。因此,探索领导对于员工与 GenAI 建议的差异感知,是理解其在二者间权衡取舍的逻辑基础。

第三,领导对于员工-GenAI 团队建议采纳的机制如何?随着团队构成形态向智能化演进,领导在吸收团队建议时面临新的挑战。传统员工团队通过社会互动与知识共享以达成共识(Bonaccio & Dalal, 2006),而 GenAI 团队则依托算法集成与多模态优化生成合成建议(Ikeda, 2024),二者在生成机制上存在本质差异,这些差异可从三个维度理解:(1)团队建议的多样性与一致性的悖论。员工团队建议天然具备异质性特征,成员间的认知差异可能催生观点碰撞或决策冲突。这种“伪多样性”可能导致领导低估 GenAI 团队建议的创新潜力,或因 GenAI 团队建议的一致性高估其可执行性(Böhm et al., 2023);(2)社会认同与算法权威的认知张力。领导评估员工团队建议时,会受到社会情感、成员地位、声望等社会线索的影响(Bonaccio & Dalal, 2006)。而 GenAI 团队缺乏社会实体属性,对 GenAI 团队的评估更多依赖技

术性能指标与系统透明性(Böhm et al., 2023)。即使 GenAI 团队建议更优,领导仍可能因“算法厌恶”而拒绝(尤其在模糊决策中)(Dietvorst et al., 2018)。但若领导理解算法原理,其态度可能转变为“算法欣赏”(Qin et al., 2025),该过程反映出技术理性与社会理性的深层冲突;(3)可追溯性与责任感的不对称。员工团队建议的决策过程可通过复盘记录、责任分配等机制进行归因(Mesmer-Magnus & DeChurch, 2009),相较之下,GenAI 因存在“黑箱”特性,其决策路径相对模糊,这意味着即便借助可解释性技术,GenAI 的决策透明度也难以媲美人类(Lundberg & Lee, 2017)。此类缺陷可能引发“算法代理悖论”:即算法的决策过程缺乏透明性归因与问责机制。这会导致在建议出错时无法厘清责任,在建议正确时又因缺乏明确的功绩归属而难以建立长期信任。基于上述原因,本研究拟探究不同类型团队建议在建议质量、风险属性及社会价值等维度上的领导差异感知。

第四,员工建议与 GenAI 建议的采纳效果差异如何?现有研究表明,GenAI 在提供建议方面与人类有显著区别,这主要体现在算法的精确度、可扩展性和个性化能力等方面(Baines et al., 2024; Kuosmanen, 2024)。从建议者特征看,GenAI 建议的透明性与可解释性是影响领导信任的关键因素。增强透明性虽有助于提升领导对 GenAI 的信任(Baines et al., 2024),但其效果受到建议准确性制约(Schmitt et al., 2021);从决策者角度看,当 GenAI 建议与领导个人预期相符时,领导更倾向于采用 GenAI 建议(Mesbah et al., 2021);从任务适配性看,GenAI 在需精细判断的任务中常不如人类建议受青睐,但随着 AI 情感理解能力的提升,其在情感类任务中的潜力正逐步显现(Daschner & Obermaier, 2022)。因此,GenAI 建议与员工建议的有效性取决于信任、透明度、任务适配性等多重维度,但现有研究并未系统探究领导对于员工-GenAI 建议采纳的效果差异(Sturm et al., 2023)。因此本研究提出如下问题:领导如何评估员工与 GenAI 建议的采纳效果?领导如何提升员工与 GenAI 建议的有效性?

第五,员工-GenAI 建议采纳的障碍是怎样的?怎样进行干预?研究表明,信任是领导采纳 GenAI 建议的核心因素,在低风险情境中尤为关键;而在高风险决策中,领导会更谨慎地评估建议的准确性与可信度(Baines et al., 2024)。同时,

感知准确性与信任形成互补关系,共同塑造领导对于 GenAI 的整体态度(Williams, 2020)。总之,领导对 GenAI 建议的采纳受到信任、感知准确性以及情境风险等多重因素的影响,但不同情境下具体采纳机制仍未明晰,故本研究提出以下问题:员工与 GenAI 建议协同的过程中面临哪些障碍?相应的干预机制是什么?以及领导应如何突破这些障碍以实现有效采纳?

2 国内外研究现状

人工智能技术的发展使得决策者可获取的建议来源日趋多元化(魏昕, 张志学, 2014; 李思贤等, 2022), GenAI 正深度融入组织决策进程。为优化决策流程、提升采纳成效,充分了解决策者对员工与 GenAI 建议的感知差异,深入探究领导者在决策过程中的心理机制显得尤为关键。

首先,现有研究对决策者的人智建议比较反应的边界条件探究不足。部分研究聚焦于决策者认知心理研究,譬如信任、算法厌恶、算法欣赏等主观因素的调节作用。具体而言:领导对 GenAI 的信任度受感知准确性影响(Williams, 2020);领导对 GenAI 的态度并非一成不变,其态度可能受到决策者心态的影响(Kim, 2026),有时表现出“算法厌恶”的决策者,在理解 GenAI 原理后也可能转变为“算法欣赏”(Qin et al., 2025)。另有研究则关注情境变量的调节效应,如决策者在数据驱动任务中更倾向于采纳 GenAI 建议(You et al., 2022),而在高风险情境中则更偏好更具伦理亲和力的员工建议(Jin & Zhang, 2025)。然而,这些研究仅聚焦于单一维度,缺乏对不同边界条件的分析,有必要对任务类型与建议策略等交互变量进行整合分析。在实际应用中,决策者常在不同任务类型中采取多样化的决策策略,而现有研究尚未充分揭示这种动态影响机制。因此,未来的研究需要更加系统地考察决策者在人智建议比较中的边界条件,构建更完善的理论框架。

其二,目前关于 GenAI 建议的研究层面存在局限。现有研究关注 GenAI 建议在组织中的实际应用,探索不同职能场景对其建议采纳的影响,相关成果主要集中于数据密集型领域,对更具普遍性的管理决策场景关注不足。在实践中,领导并非总是面临数据驱动型任务,而是需要权衡员工情感、团队协作等多维因素的复杂决策。现有

研究难以解释此类场景中领导权衡人智建议时的行为逻辑。此外,尽管部分研究提及人类-GenAI 集成决策(Choudhary et al., 2025),提出多源协同建议的效果优于单一来源,但这些探讨仅停留在效果验证层面,并未深入分析决策者的情感、认知以及行为反应。因此,有必要把 GenAI 建议采纳的研究拓展至团队层面,探讨员工、GenAI 以及员工-GenAI 团队三个层面建议之间的责任归因,以及如何影响领导的建议采纳。

其三,社会比较理论的应用维度单一,理论框架碎片化。社会比较理论在传统建议采纳研究中已被广泛引入,主要聚焦于建议者与决策者之间的能力或绩效比较对采纳行为的影响。领导在评估员工建议时,会不自觉地进行“上行比较”与“下行比较”(Matthews & Kelemen, 2025)。当领导进行上行比较时,他们可能会感到自我威胁。尤其是当下属的建议暗示其能力不足时,领导可能会为了维护自尊而拒绝该建议(Rizzo et al., 2024)。相反,平行比较和下行比较则可能使领导更倾向于采纳建议,特别是当建议被认为有助于提升组织绩效或自身能力时(Perry-Smith & Mannucci, 2017)。其他研究指出,社会比较定向是关键个体差异变量,高社会比较定向的领导对员工建议的关注度更高,易受员工表现的影响;低社会比较定向的领导则较少依赖与员工的比较,更关注建议本身的质量(Gerber et al., 2018)。这些研究为“比较心理→情感反应→采纳行为”的机制提供了支撑,但其应用场景均局限于人类建议的比较。而在人智建议采纳领域中,尽管有研究发现,高社会比较定向的领导可能为“避免人际冲突”而更倾向于采纳 GenAI 建议(Rizzo et al., 2024),但该研究仅将社会比较定向作为调节变量,未深入挖掘理论的核心维度,仍然无法解释领导对 GenAI 建议的比较逻辑。总而言之,社会比较理论应用的不足主要体现在两个层面:在比较维度上,尚未有研究探讨“社会动态比较、绩效奖赏比较、代理能力比较”在人智交互场景中的作用(Matthews & Kelemen, 2025);在理论应用层次上,也尚未阐明个体至团队的跨层次比较机制。因此有必要注重强化理论的整体性和连贯性,拓展理论应用场景,以社会比较理论作为核心框架,构建人智建议差异感知的多层次比较模型。

最后,现有研究对实践层面问题关注度较

低。现有研究多分析建议采纳的前因,例如算法技术的各种客观性能。算法准确性有助于提升建议的可信度,高准确性能够降低决策风险,进而提升领导对 GenAI 建议的采纳概率(Baines et al., 2024);可解释性作为技术属性的关键维度,领导可能因 GenAI 的“黑箱”特性对其可解释性产生质疑(Ribeiro et al., 2016),从而催生领导的不信任(Glikson & Woolley, 2020)。然而,对于如何缓解算法厌恶、降低身份威胁、管控团队建议风险等实践层面的问题,目前学界尚缺乏干预策略的探索。这极大地制约了相关理论成果转化管理实践的价值。因此,有必要分析有效干预策略,如培训、激励等,为组织优化人智协同建议采纳效能提供实践方案。

3 研究构想

本研究基于社会比较理论的社会动态、绩效奖赏与代理能力三个核心维度(Matthews & Kelemen, 2025),构建了领导对员工-GenAI 建议采纳差异化感知的整体模型。研究框架通过 5 个相互关联的子研究系统展开:研究 1 探讨了以社

会比较定向为边界条件,挖掘其对建议质量比较与欣赏程度的交互影响;研究 2 考察两种建议来源的特征差异,以相对优越性为调节变量,分析对领导情感、行为及决策有效性的作用机制;研究 3 检验两类建议在任务工作与人际工作方面的效能差异;研究 4 从团队视角,分析员工团队、GenAI 团队及员工-GenAI 团队协同建议的潜在风险,解释不同建议风险如何通过责任界定过程影响领导决策;研究 5 厘清员工-GenAI 团队协同过程中的建议障碍对建议质量的影响,以及干预策略对员工-GenAI 团队协同建议质量的提升作用。总体研究框架如图 1 所示。

3.1 研究 1: 员工-GenAI 建议质量对领导建议反应的比较研究

基于社会比较理论,本研究聚焦员工与 GenAI 的建议质量比较,探究建议质量的差异通过领导情感态度(欣赏程度)对领导采纳行为的影响,并引入社会比较定向作为边界条件,揭示领导对两类建议的差异化反应机制,构建领导对员工-GenAI 建议的情感反应过程模型。理论模型如图 2 所示。

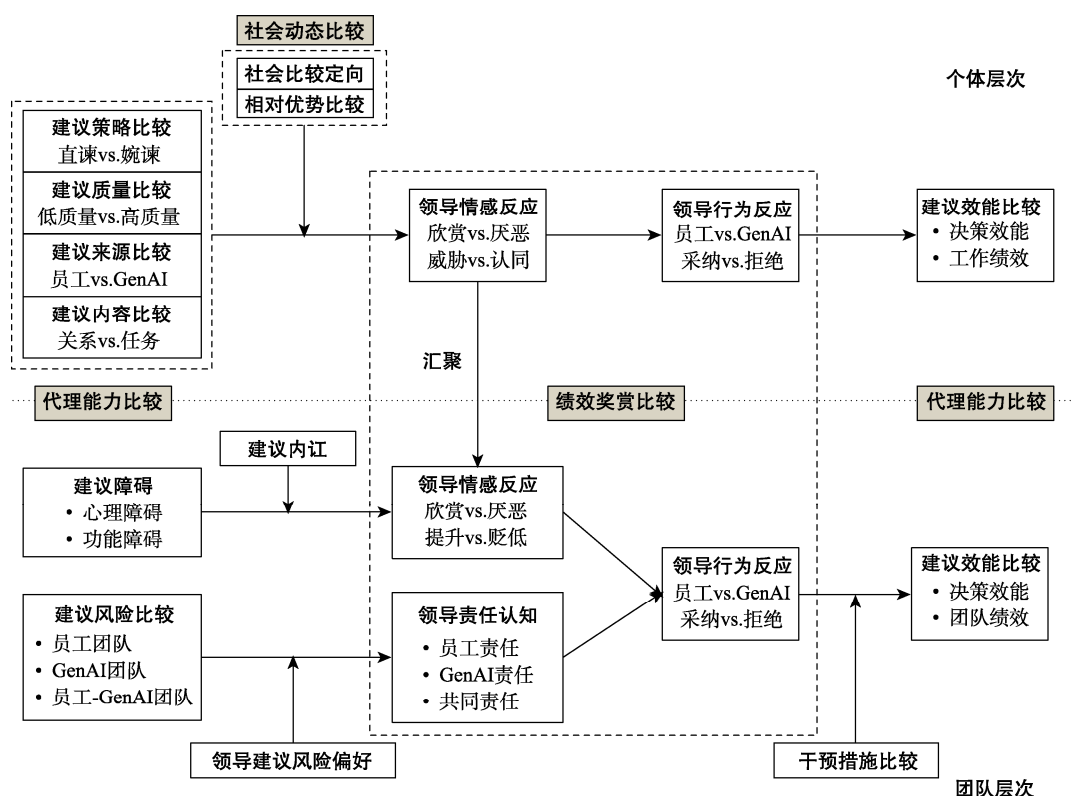


图 1 总体研究框架

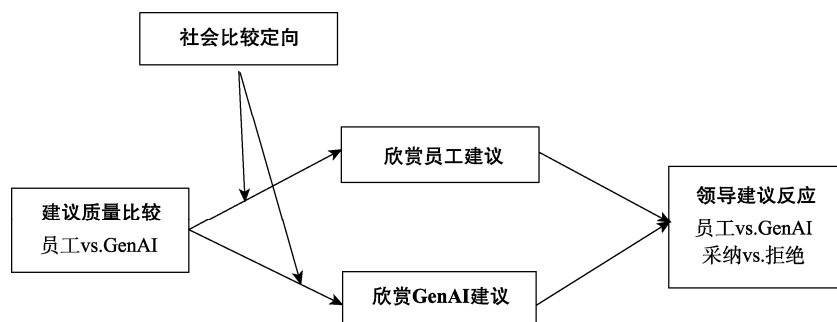


图2 员工-GenAI 建议质量的比较对领导建议反应的影响

3.1.1 建议质量比较(员工 vs. GenAI)和欣赏建议

建议质量是影响决策者采纳意愿的核心前因之一(Ikeda, 2024): 高质量的建议能够显著提高决策者采纳建议的可能性, 低质量的建议则会强化领导拒绝建议的倾向(Yaniv, 2004)。

建议质量的判断具有复杂性和多元性。现有研究对建议质量的测量主要采用 Goldsmith 和 MacGeorge (2000)、Jones 和 Bursleson (1997)、MacGeorge 等(2004)、Feng 和 MacGeorge (2010) 的量表。这类量表一般使用单维度表达, 体现建议的帮助性、适当性、敏感性、效能性和支持性(MacGeorge et al., 2002; Goldsmith & MacGeorge, 2000; MacGeorge et al., 2004; Feng & MacGeorge, 2010), 但这些研究并没有对建议质量量表进行更为细致的开发。近期, 在组织行为和人力资源管理领域一些学者开始研究谏言质量(Voice quality) (Brykman & Raver, 2021; Farh et al., 2024; Ng et al., 2022; Parke et al., 2022; Wolsink et al., 2019; Jiang et al., 2025)。学者们既有采用虚拟变量(0, 1 变量, 即谏言质量的高低) (Wolsink et al., 2019; Parke et al., 2022), 也有使用单一维度量表(Ng et al., 2022), 还有采用多维度量表(Brykman & Raver, 2021)。Brykman 和 Raver (2021)的四维度量表主要测量谏言的原则性、可行性、新颖性和组织关注。这些量表的开发为建议质量的判断提供了基础。

社会比较是一种“基本的、普遍的和强烈的人类倾向”(Corcoran et al., 2011)。基于社会比较理论(Festinger, 1954), 领导会通过对比员工与 GenAI 建议质量, 分配认知资源并形成差异化欣赏态度(Mussweiler, 2003; You et al., 2022), 换言之, 个体对不同建议的欣赏程度会随着对建议质

量认知的变化而变化(Pescetelli & Yeung, 2021)。当感知 GenAI 建议质量更优, 领导会优先关注其工具性价值, 增强欣赏程度。同时, 因认知资源有限, 领导对员工建议的欣赏则可能降低; 反之, 若领导认为员工建议质量更优, 则会提升对员工建议的欣赏, 弱化对 GenAI 建议的认可(Dang & Liu, 2024; You et al., 2022)。由此提出:

命题 1: 当领导感知 GenAI 建议质量高于员工建议质量时, 其对 GenAI 建议的欣赏程度呈正向提升(正向路径), 同时间接降低领导对员工建议的欣赏水平(负向路径)。

命题 2: 当领导感知员工建议质量高于 GenAI 建议质量时, 其对员工建议的欣赏程度呈正向提升(正向路径), 同时间接降低领导对 GenAI 建议的欣赏水平(负向路径)。

3.1.2 欣赏建议和领导建议反应

情感是连接建议评估与采纳行为的关键中介。积极情感不仅能提升个体对建议本身的喜爱度, 也能激发其帮助与慷慨行为(Ruan et al., 2024; Milyavsky & Gvili, 2024), 更能显著增强对他人的信任感(Jones & George, 1998)。因此, 积极情感会增强领导对建议者的信任, 进而提升其采纳意愿。具体而言, GenAI 建议可能被认为在技术或数学任务方面更有能力(Longoni & Cian, 2020), 而领导对 GenAI 建议的欣赏会增强其对 GenAI 能力优势的认可, 提升采纳意愿, 同时可能因对比效应降低对员工建议的接纳程度; 而对员工建议的欣赏, 会让领导更看重其贴合组织文化与情感需求的特质(Bailey et al., 2023), 倾向于采纳员工建议, 对 GenAI 建议则可能产生排斥心理(Logg et al., 2019)。

综上所述, 领导对 GenAI 建议与员工建议的

欣赏程度不仅直接影响他们对这些建议的采纳意愿,而且还因为情感状态的变化,间接地影响他们对另一类建议的接纳态度(Logg et al., 2019; You et al., 2022)。这种复杂的心理机制能够揭示领导在决策过程中如何平衡不同来源的建议,以及他们如何在欣赏与采纳之间做出权衡。由此提出:

命题 3:领导对 GenAI 建议的欣赏程度正向影响其对 GenAI 建议的采纳意愿,但负向影响对员工建议的采纳。

命题 4:领导对员工建议的欣赏程度正向影响其对员工建议的采纳意愿,但负向影响对 GenAI 建议的采纳。

3.1.3 社会比较定向的调节作用

社会比较定向指的是个体在社会情境中主动寻求或关注与他人比较的倾向性。Gibbons 和 Buunk (1999)开发了一个量表来测量社会比较定向,旨在表达自我评价在很大程度上取决于他人的表现,通常,这些个体倾向于将他人发生的事情与自己联系起来,并且对类似情况下他人的特征和成就感感兴趣。社会比较定向与目标定向相似,都表达了个体的一种稳定特质,反映了由个性所驱动的个体差异。因此,它往往被置于社会比较理论的大树之下,但它的范围更小,主要指个体的一种倾向或者动机(Gibbons & Buunk, 1999; Gerber et al., 2018)。

高社会比较定向的个体对人际关系敏感,易受他人表现影响,决策时优先关注社会评价与人际风险(Gerber et al., 2018),而低社会比较定向者则较少依赖社会比较来评估自我。这种定向差异能够调节建议质量比较与欣赏建议之间的关系:高社会比较定向领导由于对人际关系的高敏感性(Kämmer et al., 2023),会更偏好规避人际风险的选择。面对高质量员工建议,高社会比较定向领导易感知上行比较威胁(担忧自身能力被质疑)(Exline & Lobel, 1999),且会顾忌采纳特定员工的建议,以避免引发团队冲突(Rizzo et al., 2024)。而 GenAI 建议不涉及人际能力比较,既能规避自我威胁,又能避免人际冲突。因此,高社会比较定向的个体更愿意接受 GenAI 建议,而低社会比较定向的领导更专注于实际决策场景,更易欣赏员工建议的质量优势。

因此我们推断,这种个体特质也可能对不

同建议来源的情感反应和行为结果有所影响(Tai et al., 2024)。高社会比较定向的个体更在意他人的看法,也更容易受到外部参照的影响,我们推断其可能对 GenAI 建议存在偏好。反之,低社会比较定向者则偏爱员工建议。而在建议判断的行为分析中,分配给某一建议的心理权重与其采纳率直接相关(Tversky & Koehler, 1994)。由此提出:

命题 5:对于高社会比较定向的个体而言,GenAI 建议质量对其被欣赏程度的正向影响会被放大,而员工建议质量的正向影响则会被减弱。

命题 6:对于低社会比较定向的个体而言,员工建议质量对其被欣赏程度的正向影响会被放大,而 GenAI 建议质量的正向影响则会被减弱。

社会信息加工理论指出,人们对信息线索的解读具有社会情境性,同一信息在不同社会情境下可能被赋予不同的意义(Lord & Smith, 1983)。高社会比较定向的个体倾向将 GenAI 建议质量优势归因为算法系统的客观性,通过工具性的欣赏加强其采纳意向(Kaufmann, 2021);低社会比较定向者则更关注员工建议的关系属性,在员工建议质量优势时更容易产生欣赏情感,进而采纳员工建议(Chen et al., 2025)。由此提出:

命题 7:对于高社会比较定向的个体而言,GenAI 建议质量优势通过提升领导对 GenAI 建议的欣赏程度,正向影响领导采纳 GenAI 建议意愿。

命题 8:对于低社会比较定向的个体而言,员工建议质量优势通过提升领导对员工建议的欣赏程度,正向影响领导采纳员工建议意愿。

3.2 研究 2:员工-GenAI 建议来源的决策有效性比较研究

社会比较理论认为,个体在对自己和他人进行比较时,往往会产生情感反应(Dang & Liu, 2024),这种情感反应会影响其后续行为(Matthews & Kelemen, 2025)。各个来源的建议是领导进行决策的重要辅助,好的建议能够通过认知补充和效率提升来帮助领导优化决策,提高决策有效性(Chakraborty et al., 2024)。本研究对比建议源的客观特征差异,分析其通过情感中介、建议采纳/拒绝行为中介对决策有效性的影响,并引入相对优越性作为调节变量。理论模型如图 3 所示。

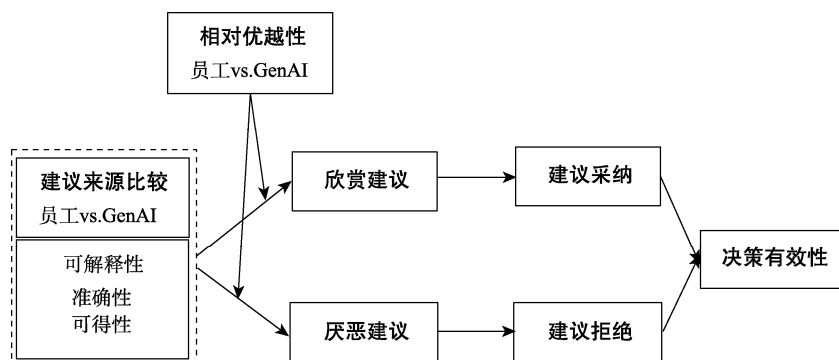


图3 员工-GenAI 建议来源的决策有效性比较

3.2.1 建议来源比较(员工 vs. GenAI)

可解释性、可得性和准确性是评价建议来源质量的重要维度,直接影响决策效能。可解释性决定了建议是否易于理解并被信任(Chakraborty et al., 2024); 准确性直接影响决策质量(Mahmud et al., 2022); 可得性影响决策效率(Brynjolfsson & McAfee, 2014)。这些维度不仅能够帮助揭示不同建议来源的优缺点,还能帮助领导在决策过程中更好地选择与依赖不同的建议来源。

可解释性是指建议的逻辑清晰度、可理解性以及因果关系的透明性(Miller, 2019)。员工建议基于个人经验与组织情境,能以自然语言阐明决策逻辑,符合人类认知框架(Endsley, 2023; Minh et al., 2022)。而 GenAI 受“黑箱”制约,底层逻辑难被非技术人员理解,还可能生成虚构数据来源等不可靠解释,因此相较于人类建议,其可解释性较弱(Wesche et al., 2022)。

准确性是指建议与客观事实、最优解或任务目标的匹配程度(Mahmud et al., 2022),其表现高度依赖任务类型与数据质量。GenAI 凭借大数据与算法优势,在数据密集型任务中通常更客观、精准(Wesche et al., 2022)。然而在模糊情境中,员工可通过社会洞察力和隐性知识提高建议准确性(Davenport & Kirby, 2016),此时 GenAI 反而可能因缺乏社会洞察力而失效。

可得性是指建议获取的便利性、响应速度和成本效率(Ramaul et al., 2024)。高可得性有助于快速响应决策需求。员工建议受时空、协调成本、组织层级限制,可能因工作负荷而延迟或信息过滤而失效(Morrison, 2011),且认知资源限制了他们处理大量数据的速度,难以快速准确地提出建

议(Luo et al., 2019)。GenAI 则能够突破时空限制,提供全天候即时服务,具有极高的可得性(Tong et al., 2021)。因此,综合来看,相较于人类员工,GenAI 建议来源的可得性更强。

3.2.2 积极反应路径: 欣赏建议和领导建议采纳

GenAI 建议所具备的高可得性和较高的准确性使其在紧急或高强度决策情境下具有显著优势,易引发领导的欣赏情感,从而促进建议采纳。具体来说,GenAI 依托海量数据和算法分析,能够有效规避人类可能的认知偏差(Glikson & Woolley, 2020),在结构化任务中其判断往往比员工建议更为准确(Tong et al., 2021)。同时,GenAI 即时响应的特性可快速满足决策需求,尤其在紧急情境下效率优势显著(Agrawal et al., 2022)。在强调效率和数据驱动的决策情境中,GenAI 能够更加高效地提供客观、无偏见的信息,这一综合优势易使领导对其产生欣赏。由于情感反应会影响其后续行为(Van Kleef, 2009),领导对 GenAI 建议的欣赏情感会进一步转化为对其的正面认知和信任,增强 GenAI 有助于解决问题、降低风险的信念,从而倾向采纳该类建议。由此提出:

命题 9: 相比于员工建议,GenAI 建议具有优势,从而导致领导欣赏 GenAI 建议,进而采纳其建议。

建议采纳的最终目的是提升决策有效性,即通过选择最优行动方案实现组织目标(Burton et al., 2020)。其核心体现在决策的客观质量与执行认可度上,如合理性、时效性等(Shamim et al., 2020)。建议来源的可解释性、准确性与可得性均与决策有效性密切相关,具备维度优势的建议更有可能优化决策成效。

在积极反应路径中, GenAI 凭借高可得性与准确性优势获得领导欣赏并被采纳。采纳此类建议能直接提升决策效率和科学性, 既能减少人为认知偏差, 增强决策客观性和准确性(Glikson & Woolley, 2020), 还能够在紧急情况下加快响应速度。因此, 采纳具备关键优势的 GenAI 建议通常对决策有效性具有积极影响(Baines et al., 2024)。由此提出:

命题 10: 建议来源会通过欣赏建议和领导建议采纳的连续中介间接地影响决策有效性。

3.2.3 消极反应路径: 厌恶建议和领导建议拒绝

尽管 GenAI 在可得性和准确性方面具有明显优势, 但“黑箱式”操作使得其可解释性较低, 易导致领导产生消极反应。由于 GenAI 决策过程往往缺乏透明性, 领导难以理解其建议的逻辑和生成机制(Burton et al., 2020; Glikson & Woolley, 2020)。相比之下, 员工的建议通常基于经验和直觉, 更容易被领导理解和接受(Doshi-Velez & Kim, 2017)。个体更倾向于信赖易于理解的信息源(Hogg et al., 1995), 因此当领导无法理解 GenAI 建议时, 易产生不信任、焦虑甚至厌恶等负面情绪(Glikson & Woolley, 2020)。在复杂或高风险的决策情境中, 领导需充分理解建议背后的逻辑以承担决策责任, 解释不足会加剧领导的不安和不信任, 从而产生厌恶情感(Ahn et al., 2021)。高情感卷入任务中, GenAI 因缺乏情感共鸣更易引发疏离与厌恶(Longoni et al., 2019; Giroux et al., 2022)。此类厌恶情绪会激发领导的损失规避心理, 促使领导为规避不确定性而拒绝 GenAI 建议(Miller, 2019)。由此提出:

命题 11: 相比于员工建议, GenAI 建议具有劣势, 从而导致领导厌恶 GenAI 建议, 进而拒绝建议。

消极反应路径中, 低可解释性使 GenAI 建议易被领导厌恶并遭到拒绝。如若 GenAI 建议客观上质量较高, 拒绝采纳会使领导错失利用高质量数据和算法支持的机会, 导致决策依赖其他信息来源, 增加主观判断风险, 降低决策准确性和效率(Cecil et al., 2024)。厌恶情绪抑制采纳意愿(Glikson & Woolley, 2020), 而领导拒绝高质量建议可能导致次优决策, 影响决策的合理性和执行效果(Damen et al., 2008)。由此提出:

命题 12: 建议来源通过厌恶建议和领导建议

拒绝的连续中介影响决策有效性。

3.2.4 相对优越性的调节作用

领导对建议来源的主观认知同样影响其反应。相对优越性是领导在比较中感知某一建议源整体优势的程度, 是关键调节变量(Choudhury & Karahanna, 2008)。该感知是对客观特征进行“心理加工”后形成的整体印象, 能放大或弱化客观特征对后续情感和行为反应的影响(Dang & Liu, 2022)。

当领导认为 GenAI 的相对优越性高时, 更可能关注其优势, 即使客观上 GenAI 的建议存在可解释性不足, 领导仍可能更倾向于欣赏并采纳其建议。优越性感知能够增强 GenAI 能力可信度(Shrestha et al., 2019), 增强决策信心(Marocco et al., 2024), 从而强化欣赏情感与采纳意向, 提升决策有效性。反之, 当领导认为 GenAI 建议来源的相对优越性低时, 更可能关注其劣势, 放大负面认知(Dietvorst et al., 2015), 强化对 GenAI 的厌恶情绪, 削弱对 GenAI 的容错意愿, 更加倾向于拒绝 GenAI 提供的建议(Glikson & Woolley, 2020)。由此提出:

命题 13: 相对优越性能够增强积极反应路径, 当领导认为 GenAI 建议来源的相对优越性高时, 更容易欣赏并采纳其建议, 从而提高决策有效性。

命题 14: 相对优越性能够削弱消极反应路径, 当领导认为 GenAI 建议来源的相对优越性低时, 更容易厌恶并拒绝其建议, 从而降低决策有效性。

3.3 研究 3: 员工-GenAI 建议内容与建议策略交互影响的比较研究

本研究基于社会比较理论, 探讨 GenAI 建议与员工建议内容(任务工作与人际工作)与建议策略(直谏与婉谏)对建议效能的影响。本研究旨在揭示领导身份威胁与身份认同在人际与任务工作中影响决策效能的中介机制, 并分析员工与 GenAI 建议在上行/下行比较中的差异性作用。理论模型如图 4 所示。

3.3.1 员工-GenAI 人际工作建议的差异化效能

社会比较理论强调, 个体倾向于通过与他人比较来评估自身能力与观点(Festinger, 1954)。在人际工作中, 员工更倾向于与具有相似背景和经验的人类建议者进行比较。人类建议者能够提供

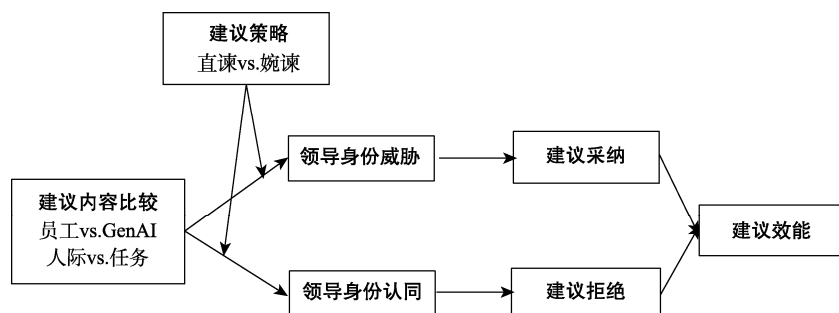


图4 员工-GenAI 建议内容与建议策略交互影响的比较研究

更具情感共鸣和情境适应性的建议,从而帮助员工更好地理解并接受建议内容,提高决策效能。同事建议作为一种亲社会行为,能使员工感知到来自同事的“善意”与支持(张若勇等, 2019)。相较之下, GenAI 缺乏情感共鸣和人际互动能力,难以提供符合员工情感需求的建议,导致决策效能较低(Baines et al., 2024)。

社会比较亦会影响个体对自身工作绩效的认知和努力程度。员工在工作中会相互观察并比较彼此的表现和成果(Matthews & Kelemen, 2025)。当员工接收到来自同事的建议时,更可能将其视为基于共同目标与经验的互助行为,从而更愿意采纳并实践该建议。此外,同事建议还能增强员工的“被关注感”,这种来自组织成员的支持会激发其以更高投入回报团队,进而提高工作绩效(张若勇等, 2019)。相反, GenAI 建议缺乏这种基于社会关系的激励作用,难以像员工建议那样显著促进绩效提升(Baines et al., 2024)。由此提出:

命题 15: 在人际工作建议中,员工建议的效能显著高于 GenAI 建议。

3.3.2 员工-GenAI 任务工作建议的差异化效能

在任务工作场景中,社会比较的重点有所不同。GenAI 具备强大的数据处理与快速分析能力,能够模拟人类的思维过程和智能行为,使机器具备独立执行生产任务的能力(Aghion et al., 2017)。在需要大量数据支持与理性分析的任务中,个体在评估建议效能时会将会将 GenAI 视为具有信息优势的参照对象。相比之下,员工建议在数据处理速度和信息全面性上可能不足,因此, GenAI 建议在此类情境中往往表现出更高的效能。GenAI 特有的深度学习和自主决策等功能,能够为员工的任务工作带来新的洞见和有价值的决策建议。对于

标准化、流程化的任务工作, GenAI 能够替代或优化重复、冗杂或低效的工作(Chui et al., 2016),其生成的建议可提供精准、规范的操作指导,有助于降低错误率,提升工作效率和质量。员工在执行任务时,其表现受个人技能水平、工作态度等因素影响,而 GenAI 建议的稳定性和准确性使其在提升工作绩效方面更具竞争力,因而 GenAI 建议比员工建议的工作绩效更高。由此提出:

命题 16: 在任务工作建议中, GenAI 建议的效能显著高于员工建议。

3.3.3 领导身份威胁和建议拒绝的中介作用

在人际工作场景中,社会比较会促使领导聚焦自身身份认知。根据社会比较理论,上行比较(与比自己优秀的人比较)可以激发自我提升的动力;而下行比较(与不如自己的人比较)可以增强自尊和信心。员工建议可能引发领导与员工之间的上行比较。若领导感知该建议暗示自身工作地位、能力受到挑战,易产生抵触情绪。有研究证实,当领导感知到自己与建议者之间存在能力或地位差距时,其对该建议的感知价值与采纳意愿会降低(段锦云, 魏秋江, 2012)。领导在感知到此类差距时,还可能会产生心理威胁或不平等感,进而影响建议采纳行为(韩翼, 肖素芳, 2020)。当员工向领导建议时,可能被解读为对领导权威的挑战或面子的冒犯,暗示员工能力优于领导,这种比较易使领导产生身份威胁。领导作为权威人物,受其面子和能力威胁的影响,会进一步倾向于拒绝员工建议。因此身份威胁和建议采纳行为在建议与建议效能间起链式中介作用。由此提出:

命题 17: 当建议内容为人际工作时,相较于员工建议, GenAI 建议更易引发领导身份威胁,进而导致领导拒绝建议,最终带来更低的建议效能。

命题 18: 当建议内容为任务工作时, 相较于 GenAI 建议, 员工建议更易引发领导身份威胁, 进而导致领导拒绝建议, 最终带来更低的建议效能。

3.3.4 领导身份认同和建议采纳的中介作用

在任务工作场景中, 领导身份认同与领导建议采纳起到了中介作用。当面对来自员工与 GenAI 的建议时, 领导者会自觉或不自觉地进行社会比较, 进而影响其对自身身份的认知与判断。

研究表明, 领导的开放性特质与员工建议采纳率密切相关。领导对建议持开放态度时其采纳建议的可能性更高(Detert et al., 2007)。裴巧玲和吴秋洁(2023)指出, 员工的建设性意见若能被领导感知, 则更易被领导采纳; 且当员工结合组织目标与领导关注点, 并以积极的、建设性的方式提出建议时, 采纳可能性会进一步增加。这种采纳行为既是对建议的积极反馈与价值肯定, 又能激励员工后续持续提出建议, 进而提升整体建议效能(章凯等, 2020)。可见, 领导身份认同感会促使建议采纳行为的发生, 二者在员工建议和建议效能之间形成链式中介作用。由此提出:

命题 19: 当建议内容为任务工作时, 相较于员工建议, GenAI 建议更易促进领导身份认同, 进而推动领导采纳建议, 最终带来更高的建议效能。

命题 20: 当建议内容为人际工作时, 相较于 GenAI 建议, 员工建议更易促进领导身份认同, 进而推动领导采纳建议, 最终带来更高的建议效能。

3.3.5 直谏策略和婉谏的差异化调节作用

在讲究人际和谐的中国社会里, 有“扬善于公堂, 规过于私室”的原则取向。即使存在不同意见, 人们也倾向于私下沟通, 表现出维护面子的行为。这表明, 不同进谏策略对领导面子威胁感知可能存在差异。例如, 直谏可能让领导感到被羞辱或冒犯(MacGeorge et al., 2004), 而婉谏对领导面子的威胁程度则相对较弱。建议策略可分为直谏和婉谏两类: 直谏是指当员工以直接的方式进谏, 其信息传递是透明的、直接的。然而, 这种直接的方式可能被管理者视为一种威胁, 因为这意味着下属在指导上级, 彰显了下属比上级知道更多、能力更强。同时, 这也宣告了上级此前的决策或方案是需要修正或错误的(Dalal &

Bonaccio, 2010)。相较之下, 婉谏是指员工以谦逊、尊重的方式提出建议, 维护了领导的尊严(Dillard et al., 1997)。采用这种策略建议时, 员工考虑到上级的感受, 从而增进领导对信息的接受程度。事实上, 婉谏也被认为是一种操控性建议策略(韩翼等, 2017), 员工采取这种方式建议顾及了领导的感受, 增强了领导可接受感。尽管直谏没有被定义为与说话者敌对的态度, 但是容易给人留下不尊重他人、武断的印象, 因此直谏可能影响领导的身份威胁感知, 而婉谏则可能影响领导的身份认同。由此提出:

命题 21: 建议内容通过领导身份威胁影响建议拒绝, 并受到直谏策略调节的影响。

命题 22: 建议内容通过领导身份认同影响建议采纳, 并受到婉谏策略调节的影响。

3.4 研究 4: 员工-GenAI 团队建议反应的责任机制比较研究

本研究基于归因理论, 分析员工团队、GenAI 团队及员工-GenAI 协同建议的潜在风险, 以及其对领导建议采纳或建议拒绝的差异化影响。为挖掘责任归因的干预机制, 研究 4 构建“建议风险—领导反应”模型, 旨在阐释不同团队建议风险如何通过责任界定过程影响领导决策, 从而揭示三种团队建议风险对领导决策产生的不同影响机制。理论模型如图 5 所示。

3.4.1 团队建议风险与领导建议反应

员工-GenAI 团队建议主要呈现三种形式: 员工团队建议、GenAI 团队建议和员工-GenAI 团队协同建议。不同类型建议存在差异化风险, 导致领导建议反应不同, 具体如下:

(1) 员工团队建议风险与领导建议反应

员工团队建议通常源于团队成员的日常观察、体验与思考, 虽然能够反映出一线经验, 但也存在一定的风险。1) 认知偏差风险: 个人经验、认知偏差和情感因素可能导致信息片面或扭曲(Kahneman et al., 2011)。2) 信息过载风险: 大量建议可能使管理层信息过载, 难以筛选有效信息(Graf & Antoni, 2023)。3) 层级压力风险: 员工建议可能因权力结构或绩效考核压力提出迎合领导偏好的非客观建议(Pfrombeck et al., 2023)。

这些风险会削弱员工建议团队的可信度, 促使领导寻求更客观的替代方案。而 GenAI 团队建议依托大数据分析 with 算法优化, 能提供系统性、

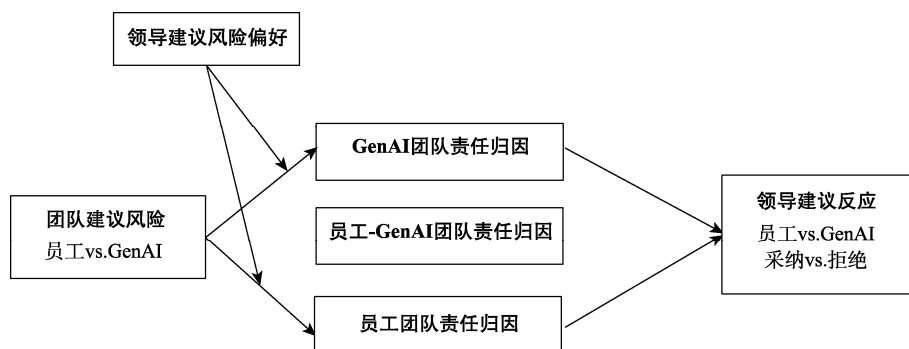


图5 员工-GenAI 建议风险及责任归因的差异性效果

客观化的决策支持(Brynjolfsson & McAfee, 2014)。员工-GenAI 协同建议则能兼顾员工创造性与GenAI 高效性,一定程度上可弥补员工团队建议的局限(Cheng et al., 2023)。由此提出:

命题 23: 员工团队建议风险会导致领导拒绝员工团队建议、采纳 GenAI 团队建议或员工-GenAI 团队协同建议。

(2) GenAI 团队建议风险与领导建议反应

GenAI 团队建议虽具备数据处理优势,但仍存在多重风险:1)算法偏见风险:输入的数据可能存在历史性偏见或不完整问题(Nelson, 2019),这可能导致不公平决策。2)过度依赖风险:长期依赖 GenAI 可能削弱员工的自主决策能力,团队创新及适应能力(Hagendorff, 2020)。若因忽视 GenAI 技术局限性引发盲目信任,还将进一步损害决策质量(Jakubik et al., 2022)。3)数据隐私与安全漏洞风险:GenAI 对大规模数据的依赖可能引发数据泄露,威胁个人隐私与企业安全。

当 GenAI 系统的建议无法满足多样化的决策需求时,领导可能会认为员工的建议更能反映实际情况和人际因素转而回归员工的经验和创造性建议。同时,员工-GenAI 协同建议的优势互补性也使其成为更优的选择(Cheng et al., 2023)。由此提出:

命题 24: GenAI 团队建议风险会导致领导拒绝 GenAI 团队建议、采纳员工团队建议或员工-GenAI 团队协同建议。

(3)员工-GenAI 协作团队建议风险与领导建议反应

员工与 GenAI 团队协同建议试图整合人智优势,但风险更复杂:1)协作冲突风险:人智认知不

一致或目标冲突(Brynjolfsson & McAfee, 2014)。

2)责任界定风险:决策失误时,难以明确责任归属(Cheng et al., 2021)。3)认知惰性风险:员工在协同决策中过度依赖 GenAI 建议导致能动性退化。这种社会惰化现象在复杂任务中尤为显著。4)伦理冲突风险:当 GenAI 建议与员工价值观冲突时,可能引发伦理争议(Siau & Wang, 2020)。

因此,员工-GenAI 协同建议风险的复杂性会使领导陷入决策困境,既无法单独信任员工建议,也难以认可 GenAI 建议,从而导致三种类型团队建议的整体拒绝。由此提出:

命题 25: 员工-GenAI 协同建议风险会导致领导拒绝员工-GenAI 三种团队建议。

3.4.2 责任归因的中介作用

基于归因理论,深入探讨了员工-GenAI 三种团队建议风险对领导建议反应的内在机制,重点分析了员工团队责任归因、GenAI 团队责任归因以及共担责任归因三个层面如何影响领导决策行为。本研究阐明了不同类型的建议风险如何通过责任归因的模式,以塑造领导对团队建议反应的过程。首先,员工团队建议风险通过员工团队责任归因影响领导的决策。当员工团队建议出现认知偏差或信息过载等问题时,领导倾向于进行内部归因,将问题归结于员工的能力局限或动机因素,这种员工团队责任归因会导致忽视员工建议(Aschauer et al., 2024),转而采纳更具客观性的 GenAI 建议或协同建议。其次,GenAI 团队建议风险通过 GenAI 责任归因影响领导的决策。面对 GenAI 建议的算法偏见或数据安全问题,领导者会进行技术归因,将责任归于系统缺陷,从而拒绝 GenAI 建议,转而寻求员工建议或协同建议中

的人类经验补充。最后,员工-GenAI 团队协同建议风险通过共担责任归因影响领导的决策。当面临协同建议的协作冲突或责任模糊时,领导者会形成共担责任归因,由于无法明确划分人智责任边界而产生决策困境,进而拒绝采纳任何一种类型的建议。

(1)员工团队责任归因的中介作用

当员工建议未能满足组织的预期或出现偏差时,管理者倾向于将失败归因于员工的个人能力或主观动机,而非考虑外部因素,这一过程被称为内在归因(Weiner, 1985)。根据归因理论,当管理者将员工建议的失败归因于员工的个人能力或情感因素时,会引起负面评价,导致员工建议被拒绝(Shaver, 2016)。与此相对,领导会将 GenAI 团队建议的算法客观性与员工-GenAI 协同建议的人智互补性视为更可靠的选择(Cheng et al., 2023)。由此提出:

命题 26: 员工团队责任归因在员工团队建议风险与领导建议反应(拒绝员工团队建议、采纳 GenAI 团队建议或员工-GenAI 团队协同建议)关系间起中介作用。

(2) GenAI 团队责任归因的中介作用

根据归因理论, GenAI 团队建议的失败往往会引发对 GenAI 系统的责任归因,特别是当其建议未能解决问题时(Miller, 2019)。GenAI 系统的建议质量可能受到数据偏见、数据不完整或算法错误的影响(Nishant et al., 2024),因此领导常将 GenAI 系统的失败归结为技术性缺陷或数据问题,而非外部环境的影响。这类归因会削弱领导对 GenAI 的信任,认为其建议缺乏可控性,进而拒绝 GenAI 团队建议;相比之下,员工团队建议的人类可控性与员工-GenAI 协同建议中人类监督的优势更易获得认可(Brynjolfsson & McAfee, 2014)。存在人类参与的建议虽然可能存在认知偏差,但相比于单纯 GenAI 系统的“黑箱”性质,员工参与的建议通常更为透明且容易被理解。基于人类与机器的对比心理,领导会优先采纳员工团队建议和员工-GenAI 团队协同建议。由此提出:

命题 27: GenAI 团队责任归因在 GenAI 团队建议风险与领导建议反应(拒绝 GenAI 团队建议、采纳员工团队建议或员工-GenAI 团队协同建议)关系间起中介作用。

(3)共担责任归因的中介作用

员工-GenAI 团队协同建议能够结合人工智能的技术优势和员工的创造力(Yue & Li, 2023),然而这种协同方式也伴随着潜在的风险,这些风险既可能源自技术缺陷,也可能来自人类操作局限。当员工-GenAI 团队协同建议的结果不理想时,领导会形成共担责任归因,即不将责任完全归因于 GenAI 或员工任何一方(Cheng et al., 2021)。

这种归因可能引发领导的情感和认知冲突,中介管理者的决策行为(Weiner, 1985)。从情绪上来说,当责任被共担时,管理者可能对 GenAI 建议和员工建议都产生不信任感,因为失败无法明确归咎于单一方(Schoenherr & Thomson, 2024)。责任边界不清可能使管理者产生对所有建议的不信任情绪,不再倾向于采纳任何一方的建议,选择回避决策风险(Jakubik et al., 2022)。最终,领导会因无法明确问责对象而拒绝员工、GenAI 及其协同团队建议。由此提出:

命题 28: 共担责任归因在员工-GenAI 团队协同建议风险与领导建议反应(拒绝员工-GenAI 三种团队建议)的关系间起中介作用。

3.4.3 领导建议风险偏好的调节作用

领导建议风险偏好,指领导评估团队建议时对潜在风险的接受容忍程度。领导对建议风险的偏好程度存在个体差异(Tversky & Kahneman, 1992),这种偏好会调节建议风险通过责任归因的中介影响领导反应的路径。风险规避型领导对责任模糊性更敏感,即使风险程度较低,也可能因为问责问题强化负面责任归因,进而拒绝建议;风险寻求型领导则更关注建议的潜在价值,可能容忍一定风险,弱化责任归因对建议采纳的负面影响,对可控风险(如局部认知偏差、可修正技术缺陷)会尝试采纳并降低风险。由此提出:

命题 29: 建议风险通过团队责任影响领导建议反应,并受到领导建议风险偏好调节的影响。

3.5 研究 5: 员工-GenAI 建议障碍与干预机制研究

基于社会比较理论,本研究探讨建议者心理与功能障碍对建议质量的影响,挖掘领导对员工和 GenAI 建议的反应差异原因,以及检验干预策略对员工和 GenAI 协同建议质量的提升作用。理论模型如图 6 所示。

3.5.1 建议障碍与建议动机

根据社会比较理论,个体通过与他人进行比

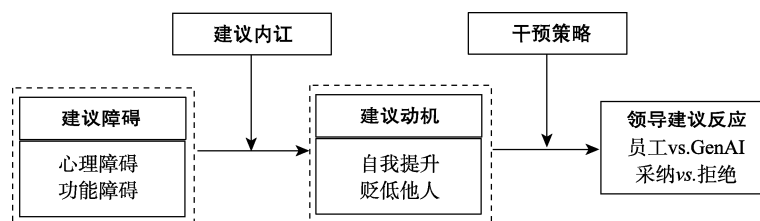


图6 员工-GenAI 建议障碍及干预策略研究

较来评估自己的能力和表现(Matthews & Kelemen, 2025)。在员工-GenAI 协同建议的场景中,员工与 GenAI 进行比较时会面临着一系列建议障碍,可以分为心理障碍和功能障碍。在宏观层面上,Rjab 等人(2023)基于技术-环境(TOE)框架将人工智能采纳障碍归纳为技术、组织与环境障碍三大类,共 18 个方面。Booyse 和 Scheepers (2024)则将引发采纳建议的障碍分为:人类社会动态、限制性法规、创造性的工作环境、缺乏信任和透明性、动态的商业环境和失去权力和控制,以及道德关怀。在微观层面上,建议障碍的因素分为两大类:功能障碍和心理障碍(Mahmud et al., 2023)。由于本文建议采纳的比较对象涉及范围是个体与团队层面,因而本研究借鉴微观层面概念,将障碍界定为心理障碍和功能障碍。功能障碍聚焦于能力技术层面,其产生于人们对采用技术所需付出的实质性改变(如使用方式、价值和风险)的感知,包括使用障碍、价值障碍和风险障碍;而心理障碍则聚焦于情感认知层面,其源于人们的既有信念与创新理念之间的冲突,包括传统障碍和印象障碍。

员工在与 GenAI 进行比较时所感知到的障碍水平,会触发员工不同的动机反应,进而影响其自身的建议行为与质量。社会比较理论认为,适度的障碍通常会激发个体的自我改善动机,通过增加努力来缩小差距(Matthews & Kelemen, 2025)。当员工意识到自己在能力上的差距时,这种比较压力往往促使他们增强努力,进而提高工作表现(Yang et al., 2025)。

然而,当员工感知到过大的障碍差距时,比较带来的压力可能导致员工焦虑感上升与自我效能感下降,反而削弱其建议动机(Edwards et al., 2024)。根据社会比较理论,当员工在面对巨大障碍时,比较带来的压力可能使他们感到自我怀疑,进而影响他们的表现(Matthews & Kelemen,

2025)。此时,员工更容易产生消极情绪和无力感,这会抑制员工建议动机,进而通过贬低 GenAI 建议获得自尊,进而降低员工建议质量(Böhm et al., 2023)。由此提出:

命题 30: 员工感知到与 GenAI 之间的适度障碍,会激发其自我提升动机,从而提升其自身建议质量。

命题 31: 员工感知到与 GenAI 之间的过大障碍,会引发其贬低他人动机,从而降低其自身建议质量。

3.5.2 建议动机的中介作用

社会比较理论指出,员工在感知自身与 GenAI 的能力差距时,可能通过提升动机以弥补差距,改善工作表现(Wesche et al., 2022)。提升动机有助于员工提出更高质量的建议,从而增加领导建议采纳的可能性(Cao et al., 2025)。因此,在员工与 GenAI 进行比较的过程中,提升动机通常能促进建议质量的改善,高质量建议进而增强领导信任,最终提升领导建议采纳概率(Rizzo et al., 2024)。

相反,员工低动机会影响建议质量,进而影响领导建议反馈。根据社会交换理论(Blau, 1964),领导对建议的反馈通常取决于员工表现。若员工因低动机而未能提升建议质量,领导更可能倾向于拒绝建议(Cao et al., 2025)。同时,社会比较理论进一步指出,当员工感知到较大障碍时,心理压力可能随之增强,进而降低建议质量,最终导致领导拒绝建议(Wood, 1996)。已有研究证明,员工对 GenAI 建议接受度较低时会产生消极情绪,进而影响建议质量(Wiesche et al., 2024)。另外,过高的 GenAI 评价标准也可能削弱员工自信,抑制建议行为,甚至影响组织创新氛围(SimanTov-Nachlieli, 2025)。在这种多重压力下,员工建议质量容易出现下降趋势,引发领导的负面反馈。由此提出:

命题 32: 建议动机在建议障碍与领导对建议的采纳之间起正向中介作用。员工因障碍引发自我提升动机,提升建议质量,从而促进领导建议采纳的可能。

命题 33: 建议动机在建议障碍与领导建议的拒绝之间起负向中介作用。员工因障碍引发贬低他人动机,导致建议质量下降,从而增加领导建议拒绝的可能。

3.5.3 建议内讧的调节作用

建议内讧指团队成员之间围绕建议内容产生的意见分歧与冲突,这一现象可能激发员工在团队工作中的竞争意识,进而影响其建议动机与行为(Bucher et al., 2024)。其核心特征可与任务冲突、团队冲突和员工建议异议等概念相关联(Kim & Cho, 2024)。研究表明,当团队成员围绕建议内容产生分歧时,这种冲突可能影响员工的建议意愿、建议质量及最终的领导反馈(Erkutlu & Chafra, 2015)。

适度的建议内讧可以视为任务冲突,即围绕建议内容的理性争论(Jehn, 1995)。适度的任务冲突有利于激发成员的思维多样性和竞争意识,从而提升建议质量,改善团队绩效(De Dreu & Weingart, 2003)。在此过程中,员工为在建议竞争中脱颖而出,可能更注重数据支撑与内容创新,进而优化建议策略(Ng & Feldman, 2012; Popelnukha et al., 2022)。

然而,若内讧超出理性范围,则可能演变为关系冲突,损害团队合作(Wibberley & Saundry, 2016)。此时员工可能会将竞争视为威胁,进而产生焦虑、不信任感,并降低建议意愿(Hyman, 2018)。此外,高度的建议冲突可能会削弱员工的情感投入,使其倾向于选择沉默或减少建设性建议(Mowbray et al., 2022)。尤其在团队内部信任水平较低的情况下,成员可能会采取防御性策略,如避免参与建议竞争、减少信息共享,甚至拒绝协作(Tangirala & Ramanujam, 2008)。

因此,建议内讧可以视为一种动态的团队建议冲突现象,其影响可能是促进性的,也可能是抑制性的。适度的建议内讧能够激发员工提升建议质量,进而提高组织绩效;但过度的建议内讧可能破坏团队信任,抑制建议行为,并影响创新能力(Kim & Cho, 2024)。由此提出:

命题 34: 建议内讧能够增强建议障碍的影响,

激发自我提升动机的效应,从而提高建议质量。

命题 35: 建议内讧削弱或逆转建议障碍的影响,引发贬低他人动机的效应,从而降低建议质量。

3.5.4 干预策略

干预策略(如培训、激励和支持)能够有效提升员工的动机和自我效能感,从而提高其工作表现(Dai et al., 2024)。研究表明,员工获得外部支持和资源有利于保持较高的工作积极性,并在面对挑战时表现出更强的适应能力(Newby et al., 2021)。有效的干预策略可通过双重路径提升员工建议质量:一方面,技能培训与心理支持能够增强员工应对障碍的能力和信心,从而提升自我效能感和建议质量(van den Heuvel et al., 2015)。另一方面,干预策略能够帮助员工适应科技变革、发挥自身优势,尤其在涉及复杂技术的环境中,能够显著增加员工高质量建议行为(Na-Nan & Sanamthong, 2020; Kim & Cho, 2024)。

然而,若干干预策略与员工实际需求不匹配,不仅无法提升信心,反而可能加剧挫败感与无力感(Dai et al., 2024)。当员工感知自身与 GenAI 技术的差距过大,而组织培训又无法有效弥补时,其建议动机将显著降低,甚至产生消极情绪。根据社会比较理论(Festinger, 1954),无效干预会降低自我效能感,从而抑制建议行为(Xavier & Korunka, 2025)。具体而言,脱离实际需求的培训内容、缺乏长期跟踪和反馈机制,或不当的激励机制,都有可能让员工感受到自己的努力未被认可,导致工作动机下降(Wang & Chuang, 2024),降低其建议意愿(Kim et al., 2020)。

因此,干预策略在员工建议中的作用取决于其实施效果。有效的培训和激励措施能够增强员工的自我效能感,提高工作动机,并提升建议质量(Schemmer et al., 2023)。而无效干预则可能适得其反,抑制建议行为(Newby et al., 2021)。因此,管理者需根据员工实际情况提供个性化支持,以确保干预策略发挥积极作用(van den Heuvel et al., 2015)。由此提出:

命题 36: 有效的干预策略会增强正向激励效应,进而增强自我提升动机对领导建议采纳的影响。

命题 37: 无效的干预策略会削弱正向激励效应,甚至产生负向影响,进而增强贬低他人动机对领导建议拒绝的影响。

4 理论建构

本文基于社会比较理论,从社会动态比较、绩效奖赏比较和代理能力比较三个方面,系统揭示组织管理情境中领导对员工与 GenAI 建议采纳的差异感知机制,构建多层次理论框架,旨在探索领导对员工—GenAI 建议采纳的认知反应、情感反应、行为反应及动态影响机制,为组织决策优化、人智协同效能提升及人工智能治理提供科学依据。本文理论贡献主要体现在以下四个方面:

第一,本文拓展了社会比较理论的应用边界与维度体系,搭建多维度人智比较框架,实现从人类间比较到人智比较的理论突破。传统的社会比较理论多聚焦于人与人之间的比较,但 GenAI 的技术突破为该理论提出了非人类比较对象的新情境,因此本文将社会比较理论拓展至人智比较场景,整合了社会动态比较、绩效奖赏比较以及代理能力比较的比较框架,厘清该框架在领导对 GenAI 建议反应中的关键作用。通过引入建议来源、建议特征、建议内容、建议质量、建议采纳风险和障碍等多方面的比较,本文不仅丰富了社会比较的维度,也揭示了人智比较中特有的认知与情感机制,为社会比较理论在人工智能时代的演进提供了新的理论路径。

第二,本文构建了个体到团队的跨层次领导人智建议反应整合模型。本研究从个体层面的建议质量感知、情感反应与行为采纳,延伸至团队层面的责任归因、风险感知与干预机制,形成了从微观认知到宏观决策的完整理论链条。通过引入责任归因与建议内证等团队层面的中介与调节变量,本文揭示了团队建议风险影响领导决策的内在机制,弥补现有研究在团队层面人智协同建议采纳机制上的理论空白,更丰富了社会比较理论在团队层面上的研究层次。

第三,提出了 GenAI 建议采纳的双刃剑效应的理论框架并验证其干预策略。本文系统识别 GenAI 建议在提升决策效率与客观性的同时,也可能引发身份威胁、责任模糊与信任危机等负面效应。通过引入建议障碍与干预策略,本文不仅揭示了建议障碍对建议动机与质量的影响机制,还构建了干预路径模型,这能够帮助企业明确从哪些层面干预人智建议采纳,也为未来关于人智团队建议采纳干预策略的研究提供路径参考与

帮助。

第四,推动了管理学、心理学与人工智能研究的跨学科融合。本文整合了社会比较理论、归因理论、建议采纳理论与人工智能信任研究,构建了多理论交叉的研究框架。通过将可解释性、可得性、准确性等 GenAI 技术属性与社会比较定向、身份威胁、责任归因等心理与社会变量相结合,本文不仅在理论上拓展了人智交互研究的深度,也在方法上为后续跨学科研究提供了可操作的变量体系与模型范式。

参考文献

- 段锦云,魏秋江.(2012).建言效能感结构及其在员工建言行为发生中的作用.《心理学报》,44(7),972-985.
- 韩翼,肖素芳.(2020).领导为什么拒谏:基于动机社会认知视角的阐释.《外国经济与管理》,42(8),68-80.
- 韩翼,董越,胡筱菲,谢怡敏.(2017).员工进谏策略及其有效性研究.《管理学报》,14(12),1777-1785.
- 李思贤,陈佳昕,宋艾珈,王梦琳,段锦云.(2022).人们对人工智能建议接受度的影响因素.《心理技术与应用》,10(4),202-214.
- 刘伟,谭文辉.(2025).人机环境系统融合智能:超越人类智能的可能性.清华大学出版社.
- 裴巧玲,吴秋洁.(2023).员工建言与管理者纳言:一个有调节的中介模型.《现代营销》,(12),119-121.
- 魏昕,张志学.(2014).上级何时采纳促进性或抑制性进言?上级地位和下属专业度的影响.《管理世界》,(1),132-143.
- 许丽颖,赵一骏,喻丰.(2025).人工智能主管提出的道德行为建议更少被遵从.《心理学报》,57(11),2060-2082.
- 章凯,时金京,罗文豪.(2020).建言采纳如何促进员工建言:基于目标自组织视角的整合机制.《心理学报》,52(2),229-239.
- 张若勇,闫石,邵琪.(2019).同事建言对员工任务绩效影响机制的研究.《兰州大学学报(社会科学版)》,47(4),73-82.
- 宗树伟,杨付,龙立荣,韩翼.(2025).促进还是抑制?生成式人工智能建议采纳对创造力的双刃剑效应.《心理科学进展》,33(6),905-915.
- Aghion, P., Jones, B. F., & Jones, C. I. (2017). Artificial intelligence and economic growth. In Agrawal, A., Gans, J., & Goldfarb, A. (Eds.), *National bureau of economic research* (pp.237-290). University of Chicago Press.
- Agrawal, A., Gans, J., & Goldfarb, A. (2022). *Power and prediction: The disruptive economics of artificial intelligence*. Harvard Business Review Press.
- Agrawal, A., Gans, J. S., & Goldfarb, A. (2024). Artificial intelligence adoption and system-wide change. *Journal of Economics & Management Strategy*, 33(2), 327-337.
- Ahn, P., Swol, L., Kim, S. J., & Park, H. (2021). Enhanced motivation and decision making from going hybrid. *Small Group Research*, 53, 104649642110435.

- Aschauer, F., Sohn, M., & Hirsch, B. (2024). Managerial advice-taking-Sharing responsibility with (non) human advisors trumps decision accuracy. *European Management Review*, 21(1), 186–203.
- Bailey, P. E., Leon, T., Ebner, N. C., Moustafa, A. A., & Weidemann, G. (2023). A meta-analysis of the weight of advice in decision-making. *Current Psychology*, 42(28), 24516–24541.
- Baines, J. I., Dalal, R. S., Ponce, L. P., & Tsai, H. C. (2024). Advice from artificial intelligence: A review and practical implications. *Frontiers in Psychology*, 15, 1390182.
- Blau, P. M. (1964). *Exchange and power in social life*. New York, NY: Wiley.
- Bonaccio, S., & Dalal, R. S. (2006). Advice taking and decision-making: An integrative literature review, and implications for the organizational sciences. *Organizational Behavior and Human Decision Processes*, 101(2), 127–151.
- Booyse, D., & Scheepers, C. B. (2024). Barriers to adopting automated organisational decision-making through the use of artificial intelligence. *Management Research Review*, 47(1), 64–85.
- Böhm, R., Jörling, M., Reiter, L., & Fuchs, C. (2023). People devalue generative AI's competence but not its advice in addressing societal and personal challenges. *Communications Psychology*, 1(1), 1–10.
- Brykman, K. M., & Raver, J. L. (2021). To speak up effectively or often? The effects of voice quality and voice frequency on peers' and managers' evaluations. *Journal of Organizational Behavior*, 42(4), 504–526.
- Brynjolfsson, E., & McAfee, A. (2014). *The second machine age: Work, progress, and prosperity in a time of brilliant technologies*. W.W. Norton & Company.
- Bucher, E., Schou, P. K., & Waldkirch, M. (2024). Just another voice in the crowd? Investigating digital voice formation in the gig economy. *Academy of Management Discoveries*, 10(3), 488–511.
- Burton, J. W., Stein, M.-K., & Jensen, T. B. (2020). A systematic review of algorithm aversion in augmented decision making. *Journal of Behavioral Decision Making*, 33(2), 220–239.
- Cao, X., De Zwaan, L., & Wong, V. (2025). Building trust in robo-advisory: Technology, firm-specific and system trust. *Qualitative Research in Financial Markets*. <https://doi.org/10.1108/QRFM-02-2024-0033>
- Cecil, J., Lermer, E., Hudecek, M. F., Sauer, J., & Gaube, S. (2024). Explainability does not mitigate the negative impact of incorrect AI advice in a personnel selection task. *Scientific Reports*, 14(1), 9736. <https://doi.org/10.1038/s41598-024-60220-5>
- Chakraborty, D., Kar, A. K., Patre, S., & Gupta, S. (2024). Enhancing trust in online grocery shopping through generative AI chatbots. *Journal of Business Research*, 180, 114–137.
- Chen, Q., Yin, C., & Gong, Y. (2025). Would an AI chatbot persuade you: An empirical answer from the elaboration likelihood model. *Information Technology & People*, 38(2), 937–962.
- Cheng, L., Varshney, K. R., & Liu, H. (2021). Socially responsible AI algorithms: Issues, purposes, and challenges. *Journal of Artificial Intelligence Research*, 71, 1137–1181.
- Cheng, B., Lin, H., & Kong, Y. (2023). Challenge or hindrance? How and when organizational artificial intelligence adoption influences employee job crafting. *Journal of Business Research*, 164, 113987. <https://doi.org/10.1016/j.jbusres.2023.113987>
- Chui, M., Manyika, J., & Miremadi, M. (2016). Where machines could replace humans-and where they can't (yet). *The McKinsey Quarterly*, (3), 58–69.
- Choudhury, V., & Karahanna, E. (2008). The relative advantage of electronic channels: A multidimensional view. *MIS Quarterly*, 32(1), 179–200.
- Choudhary, V., Marchetti, A., Shrestha, Y. R., & Puranam, P. (2025). Human-AI ensembles: When can they work?. *Journal of Management*, 51(2), 536–569.
- Civit, M., Civit-Masot, J., Cuadrado, F., & Escalona, M. J. (2022). A systematic review of artificial intelligence-based music generation: Scope, applications, and future trends. *Expert Systems with Applications*, 209, 118190. <https://doi.org/10.1016/j.eswa.2022.118190>
- Corcoran, K., Crusius, J., & Mussweiler, T. (2011). Social comparison: Motives, standards, and mechanisms. In D. Chadee (Ed.), *Theories in Social Psychology* (pp. 119–139). Wiley Blackwell.
- Dai, Y., Lee, J., & Kim, J. W. (2024). AI vs. human voices: How delivery source and narrative format influence the effectiveness of persuasion messages. *International Journal of Human-Computer Interaction*, 40(24), 8735–8749.
- Dalal, R. S., & Bonaccio, S. (2010). What types of advice do decision-makers prefer? *Organizational Behavior and Human Decision Processes*, 112(1), 11–23.
- Damen, F., Van Knippenberg, B., & Van Knippenberg, D. (2008). Affective match in leadership: Leader emotional displays, follower positive affect, and follower performance. *Journal of Applied Social Psychology*, 38(4), 868–902.
- Dang, J., & Liu, L. (2022). Implicit theories of the human mind predict competitive and cooperative responses to AI robots. *Computers in Human Behavior*, 134, 107300. <https://doi.org/10.1016/j.chb.2022.107300>
- Dang, J., & Liu, L. (2024). Extended artificial intelligence aversion: People deny humanness to artificial intelligence users. *Journal of Personality and Social Psychology*. <https://doi.org/10.1037/pspi0000480>
- Daschner, S., & Obermaier, R. (2022). Algorithm aversion? On the influence of advice accuracy on trust in algorithmic advice. *Journal of Decision Systems*, 31(sup1), 77–97.
- Davenport, T. H., & Kirby, J. (2016). *Only humans need*

- apply: *Winners and losers in the age of smart machines*. Harper Business.
- De Dreu, C. K. W., & Weingart, L. R. (2003). Task versus relationship conflict, team performance, and team member satisfaction: A meta-analysis. *Journal of Applied Psychology*, 88(4), 741–749.
- Detert, J. R., & Burris, E. R. (2007). Leadership behavior and employee voice: Is the door really open? *Academy of Management Journal*, 50(4), 869–884.
- Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion: People erroneously avoid algorithms after seeing them err. *Journal of Experimental Psychology: General*, 144(1), 114–126.
- Dietvorst, B. J., Simmons, J. P., & Massey, C. (2018). Overcoming algorithm aversion: People will use imperfect algorithms if they can (even slightly) modify them. *Management Science*, 64(3), 1155–1170.
- Dillard, J. P., Wilson, S. R., Tusing, K. J., & Kinney, T. A. (1997). Politeness judgments in personal relationships. *Journal of Language and Social Psychology*, 16(3), 297–325.
- Doshi-Velez, F., & Kim, B. (2017). Towards a rigorous science of interpretable machine learning. *PsyArXiv*. <https://doi.org/10.48550/arXiv.1702.08608>
- Edwards, M. R., Zubielevitch, E., Okimoto, T., Parker, S., & Anseel, F. (2024). Managerial control or feedback provision: How perceptions of algorithmic HR systems shape employee motivation, behavior, and well-being. *Human Resource Management*, 63(4), 691–710.
- Eidsley, M. R. (2023). Supporting human-ai teams: Transparency, explainability, and situation awareness. *Computers in Human Behavior*, 140, 107574. <https://doi.org/10.1016/j.chb.2022.107574>
- Erkutlu, H., & Chafra, J. (2015). The mediating roles of psychological safety and employee voice on the relationship between conflict management styles and organizational identification. *American Journal of Business*, 30(1), 72–91.
- Exline, J. J., & Lobel, M. (1999). The perils of outperformance: Sensitivity about being the target of a threatening upward comparison. *Psychological Bulletin*, 125(3), 307–337.
- Farh, C. I., Li, J., & Lee, T. W. (2024). Toward a contextualized view of voice quality, its dimensions, and its dynamics across newcomer socialization. *Academy of Management Review*, 49(2), 399–428.
- Feng, B., & MacGeorge, E. L. (2010). The influences of message and source factors on advice outcomes. *Communication Research*, 37(4), 553–575.
- Festinger, L. (1954). A theory of social comparison processes. *Human Relations*, 7(2), 117–140.
- Feuerriegel, S., Hartmann, J., Janiesch, C., & Zschech, P. (2024). Generative AI. *Business & Information Systems Engineering*, 66(1), 111–126.
- Gerber, J. P., Wheeler, L., & Suls, J. (2018). A social comparison theory meta-analysis 60+ years on. *Psychological Bulletin*, 144(2), 177–197.
- Gibbons, F. X., & Buunk, B. P. (1999). Individual differences in social comparison: Development of a scale of social comparison orientation. *Journal of Personality and Social Psychology*, 76(1), 129–142.
- Giroux, M., Kim, J., Lee, J. C., & Park, J. (2022). Artificial intelligence and declined guilt: Retailing morality comparison between human and AI. *Journal of Business Ethics*, 178(4), 1027–1041.
- Glikson, E., & Woolley, A. W. (2020). Human trust in artificial intelligence: Review of empirical research. *Academy of Management Annals*, 14(2), 627–660.
- Goldsmith, D. J., & MacGeorge, E. L. (2000). The impact of politeness and relationship on perceived quality of advice about a problem. *Human Communication Research*, 26, 234–263.
- Graf, B., & Antoni, C. H. (2023). Drowning in the flood of information: A meta-analysis on the relation between information overload, behaviour, experience, and health and moderating factors. *European Journal of Work and Organizational Psychology*, 32(2), 173–198.
- Hagendorff, T. (2020). The ethics of AI ethics: An evaluation of guidelines. *Minds and Machines*, 30(1), 99–120.
- Hogg, M. A., Terry, D. J., & White, K. M. (1995). A tale of two theories: A critical comparison of identity theory with social identity theory. *Social Psychology Quarterly*, 58(4), 255–269.
- Huang, Z., Che, C., Zheng, H., & Li, C. (2024). Research on generative artificial intelligence for virtual financial robo-advisor. *Academic Journal of Science and Technology*, 10(1), 74–80.
- Hyman, J. (2018). *Employee voice and participation: Contested past, troubled present, uncertain future*. Routledge.
- Ikeda, S. (2024). Inconsistent advice by ChatGPT influences decision making in various areas. *Scientific Reports*, 14(1), 15876. <https://doi.org/10.1038/s41598-024-66821-4>
- Jakubik, J., Schöffner, J., Hoge, V., Vössing, M., & Kühl, N. (2022). An empirical evaluation of predicted outcomes as explanations in human-AI decision-making. In M. R., Amini, S., Canu, A., Fischer, T., Guns, P. K., Novak & G., Tsoumakas (Eds), *Selected paper on machine learning and knowledge discovery in databases* (pp. 353–368). Springer-Verlag.
- Jeblick, K., Schachtner, B., Dexl, J., Mittermeier, A., Stüber, A. T., Topalis, J., ... Ingris, M. (2024). ChatGPT makes medicine easy to swallow: An exploratory case study on simplified radiology reports. *European Radiology*, 34(5), 2817–2825.
- Jehn, K. A. (1995). A multimethod examination of the benefits and detriments of intragroup conflict. *Administrative Science Quarterly*, 40(2), 256–282.
- Jiang, J., Gong, Y., Dong, Y., Han, Y., & Qin, Y. (2025).

- Unpacking leader critical thinking in employee voice quality and silence frequency. *Journal of Occupational and Organizational Psychology*, 98(1), e12554. <https://doi.org/10.1111/joop.12554>
- Jiang, H. H., Brown, L., Cheng, J., Khan, M., Gupta, A., Workman, D., ... Gebru, T. (2023). AI art and its impact on artists. In F., Rossi, S., Das, J., Davis, K., Firth-Butterfield, & A., John (Eds), *Selected paper on proceedings of the 2023 AAAI/ACM conference on AI, ethics, and society* (pp. 363–374). Association for Computing Machinery.
- Jin, F., & Zhang, X. (2025). Artificial intelligence or human: When and why consumers prefer AI recommendations. *Information Technology & People*, 38(1), 279–303.
- Jones, G. R., & George, J. M. (1998). The experience and evolution of trust: Implications for cooperation and teamwork. *Academy of Management Review*, 23(3), 531–546.
- Jones, S. M., & Burleson, B. R. (1997). The impact of situational variables on helpers' perceptions of comforting messages: An attributional analysis. *Communication Research*, 24(5), 530–555.
- Kahneman, D., Lovallo, D., & Sibony, O. (2011). Before you make that big decision. *Harvard Business Review*, 89(6), 50–60.
- Kahr, P. K., Rooks, G., Snijders, C., & Willemsen, M. C. (2024). The trust recovery journey. The effect of timing of errors on the willingness to follow AI advice. In B., Knijnenburg (Eds), *Selected paper on proceedings of the 29th international conference on intelligent user interfaces* (pp. 609–622). Association for Computing Machinery.
- Kämmer, J. E., Choshen-Hillel, S., Müller-Trede, J., Black, S. L., & Weibler, J. (2023). A systematic review of empirical studies on advice-based decisions in behavioral and organizational research. *Decision*, 10(2), 107–137.
- Kasneci, E., Sessler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., ... Kasneci, G. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103, 102274. <https://doi.org/10.1016/j.lindif.2023.102274>
- Kaufmann, E. (2021). Algorithm appreciation or aversion? Comparing in-service and pre-service teachers' acceptance of computerized expert models. *Computers and Education: Artificial Intelligence*, 2, 100028.
- Kim, H. (2026). Beyond algorithm aversion: The impact of psychological readiness on algorithmic advice. *Computers in Human Behavior*, 174, 108824. <https://doi.org/10.1016/j.chb.2025.108824>
- Kim, H. Y., Lee, Y. S., & Jun, D. B. (2020). Individual and group advice taking in judgmental forecasting: Is group forecasting superior to individual forecasting? *Journal of Behavioral Decision Making*, 33(3), 287–303.
- Kim, T., & Cho, W. (2024). Employee voice opportunities enhance organizational performance when faced with competing demands. *Review of Public Personnel Administration*, 44(4), 713–739.
- Kuosmanen, O. J. (2024). *Advice from humans and artificial intelligence: Can we distinguish them, and is one better than the other?* [Unpublished master's thesis]. UiT Norges arktiske universitet.
- Logg, J. M., Minson, J. A., & Moore, D. A. (2019). Algorithm appreciation: People prefer algorithmic to human judgment. *Organizational Behavior and Human Decision Processes*, 151, 90–103.
- Longoni, C., Bonezzi, A., & Morewedge, C. K. (2019). Resistance to medical artificial intelligence. *Journal of Consumer Research*, 46(4), 629–650.
- Longoni, C., & Cian, L. (2020). Artificial intelligence in utilitarian vs. Hedonic contexts: The “word-of-machine” effect. *Journal of Marketing*, 86(1), 96–118.
- Lord, R. G., & Smith, J. E. (1983). Theoretical, information processing, and situational factors affecting attribution theory models of organizational behavior. *Academy of Management Review*, 8(1), 50–60.
- Lundberg, S. M., & Lee, S.-I. (2017). A unified approach to interpreting model predictions. *Proceedings of the 31st International Conference on Neural Information Processing Systems*, 4768–4777.
- Luo, X., Tong, S., Fang, Z., & Qu, Z. (2019). Frontiers: Machines vs. Humans: the impact of artificial intelligence chatbot disclosure on customer purchases. *Marketing Science*, 38(6), 937–947.
- MacGeorge, E. L., Feng, B., Butler, G. L., & Budarz, S. K. (2004). Understanding advice in supportive interactions: Beyond the facework and message evaluation paradigm. *Human Communication Research*, 30(1), 42–70.
- MacGeorge, E. L., Lichtman, R. M., & Pressey, L. C. (2002). The evaluation of advice in supportive interactions: Facework and contextual factors. *Human Communication Research*, 28(3), 451–463.
- Mahmud, H., Islam, A. N., Ahmed, S. I., & Smolander, K. (2022). What influences algorithmic decision-making? A systematic literature review on algorithm aversion. *Technological Forecasting and Social Change*, 175, 121390.
- Mahmud, H., Islam, A. N., & Mitra, R. K. (2023). What drives managers towards algorithm aversion and how to overcome it? Mitigating the impact of innovation resistance through technology readiness. *Technological Forecasting and Social Change*, 193, 122641. <https://doi.org/10.1016/j.techfore.2023.122641>
- Marocco, S., Talamo, A., & Quintiliani, F. (2024). From service design thinking to the third generation of activity theory: A new model for designing AI-based decision-support systems. *Frontiers in Artificial Intelligence*, 7, 1303691. <https://doi.org/10.3389/frai.2024.1303691>

- Matthews, M. J., & Kelemen, T. K. (2025). To compare is human: A review of social comparison theory in organizational settings. *Journal of Management*, 51(1), 212–248.
- Mesbah, N., Tauchert, C., & Buxmann, P. (2021). Whose advice counts more—man or machine? An experimental investigation of ai-based advice utilization. *Proceedings of the 54th Hawaii International Conference on System Sciences* (pp.4083–4092). Kauai, Hawaii, USA.
- Mesmer-Magnus, J. R., & DeChurch, L. A. (2009). Information sharing and team performance: A meta-analysis. *Journal of Applied Psychology*, 94(2), 535–546.
- Miller, T. (2019). Explanation in artificial intelligence: Insights from the social sciences. *Artificial Intelligence*, 267, 1–38.
- Milyavsky, M., & Gvili, Y. (2024). Advice taking vs. combining opinions: Framing social information as advice increases source's perceived helping intentions, trust, and influence. *Organizational Behavior and Human Decision Processes*, 183, 104328. <https://doi.org/10.1016/j.obhdp.2024.104328>
- Minh, D., Wang, H. X., Li, Y. F., & Nguyen, T. N. (2022). Explainable artificial intelligence: A comprehensive review. *Artificial Intelligence Review*, 55(5), 3503–3568.
- Morrison, E. W. (2011). Employee voice behavior: Integration and directions for future research. *Academy of Management Annals*, 5(1), 373–412.
- Mowbray, P. K., Wilkinson, A., & Tse, H. H. (2022). Strategic or silencing? Line managers' repurposing of employee voice mechanisms for high performance. *British Journal of Management*, 33(2), 1054–1070.
- Mussweiler, T. (2003). Comparison processes in social judgment: Mechanisms and consequences. *Psychological Review*, 110(3), 472–489.
- Na-Nan, K., & Sanamthong, E. (2020). Self-efficacy and employee job performance: Mediating effects of perceived workplace support, motivation to transfer and transfer of training. *International Journal of Quality & Reliability Management*, 37(1), 1–17.
- Nelson, G. S. (2019). Bias in artificial intelligence. *North Carolina Medical Journal*, 80(4), 220–222.
- Newby, J., Mason, E., Kladnitski, N., Murphy, M., Millard, M., Haskelberg, H., ... Mahoney, A. (2021). Integrating internet CBT into clinical practice: A practical guide for clinicians. *Clinical Psychologist*, 25(2), 164–178.
- Ng, T. W., & Feldman, D. C. (2012). Employee voice behavior: A meta-analytic test of the conservation of resources framework. *Journal of Organizational Behavior*, 33(2), 216–234.
- Ng, T. W., Wang, M., Hsu, D. Y., & Su, C. (2022). Voice quality and ostracism. *Journal of Management*, 48(2), 281–318.
- Nishant, R., Schneckenberg, D., & Ravishankar, M. N. (2024). The formal rationality of artificial intelligence-based algorithms and the problem of bias. *Journal of Information Technology*, 39(1), 19–40.
- Oksanen, A., Cvetkovic, A., Akin, N., Latikka, R., Bergdahl, J., Chen, Y., & Savela, N. (2023). Artificial intelligence in fine arts: A systematic review of empirical research. *Computers in Human Behavior: Artificial Humans*, 1(2), 100004. <https://doi.org/10.1016/j.chbah.2023.100004>
- Parke, M. R., Tangirala, S., Sanaria, A., & Ekkirala, S. (2022). How strategic silence enables employee voice to be valued and rewarded. *Organizational Behavior and Human Decision Processes*, 173, 104187. <https://doi.org/10.1016/j.obhdp.2022.104187>
- Perry-Smith, J. E., & Mannucci, P. V. (2017). From creativity to innovation: The social network drivers of the four phases of the idea journey. *Academy of Management Review*, 42(1), 53–79.
- Pescetelli, N., & Yeung, N. (2021). The role of decision confidence in advice-taking and trust formation. *Journal of Experimental Psychology: General*, 150(3), 507–520.
- Pfrombeck, J., Levin, C., Rucker, D. D., & Galinsky, A. D. (2023). The hierarchy of voice framework: The dynamic relationship between employee voice and social hierarchy. *Research in Organizational Behavior*, 100179. <https://doi.org/10.1016/j.riob.2022.100179>
- Popelnukha, A., Almeida, S., Obaid, A., Sarwar, N., Atamba, C., Tariq, H., & Weng, Q. (2022). Keep your mouth shut until I feel good: Testing the moderated mediation model of leader's threat to competence, self-defense tactics, and voice rejection. *Personnel Review*, 51(1), 394–431.
- Qin, X., Zhou, X., Chen, C., Wu, D., Zhou, H., Dong, X., ... Lu, J. G. (2025). AI aversion or appreciation? A capability-personalization framework and a meta-analytic review. *Psychological Bulletin*, 151(5), 580–599.
- Ramaul, L., Ritala, P., & Ruokonen, M. (2024). Creational and conversational AI affordances: How the new breed of chatbots is revolutionizing knowledge industries. *Business Horizons*, 67(5), 615–627.
- Rjab, A. B., Mellouli, S., & Corbett, J. (2023). Barriers to artificial intelligence adoption in smart cities: A systematic literature review and research agenda. *Government Information Quarterly*, 40(3), 101814. <https://doi.org/10.1016/j.giq.2023.101814>
- Ribeiro, M. T., Singh, S., & Guestin, C. (2016). “Why should I trust you?”: Explaining the predictions of any classifier. *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp. 1135–1144). San Diego, California, USA.
- Rizzo, C., Bagna, G., & Tuček, D. (2024). Do managers trust AI? An exploratory research based on social comparison theory. *Management Decision*, 63(10), 3625–3641. <https://doi.org/10.1108/MD-10-2023-1971>
- Ruan, Y., Le, J. D. V., & Reis, H. T. (2024). How can I help?:

- Specific strategies used in interpersonal emotion regulation in a relationship context. *Emotion*, 24(2), 329–344.
- Sachin, P. K., & Schecter, A. (2024). Advice utilization in combined human-algorithm decision-making: An analysis of preferences and behaviors. *Journal of the Association for Information Systems*, 25(6), 1439–1465.
- Schemmer, M., Bartos, A., Spitzer, P., Hemmer, P., Kühl, N., Liebschner, J., & Satzger, G. (2023). Towards effective human-AI decision-making: The role of human learning in appropriate reliance on AI advice. *PsyArXiv*. <https://doi.org/10.48550/arXiv.2310.02108>
- Schmitt, A., Zierau, N., Janson, A., & Leimeister, J. M. (2021, June). *Voice as a contemporary frontier of interaction design*. Paper presented at the meeting of European conference on information systems 2021 Research Papers. Marrakech, Morocco.
- Schoenherr, J. R., & Thomson, R. (2024). When AI fails, who do we blame? Attributing responsibility in human-AI interactions. *IEEE Transactions on Technology and Society*, 5(1), 61–70.
- Shamim, S., Zeng, J., Khan, Z., & Zia, N. U. (2020). Big data analytics capability and decision making performance in emerging market firms: The role of contractual and relational governance mechanisms. *Technological Forecasting and Social Change*, 161, 120315. <https://doi.org/10.1016/j.techfore.2020.120315>
- Shaver, K. G. (2016). *Introduction to attribution processes*. Routledge.
- Shrestha, Y. R., Ben-Menahem, S. M., & Von Krogh, G. (2019). Organizational decision-making structures in the age of artificial intelligence. *California Management Review*, 61(4), 66–83.
- Siau, K., & Wang, W. (2020). Artificial intelligence (AI) ethics: Ethics of AI and ethical AI. *Journal of Database Management*, 31(2), 74–87.
- SimanTov-Nachlieli, I. (2025). More to lose: The adverse effect of high performance ranking on employees' preimplementation attitudes toward the integration of powerful AI aids. *Organization Science*, 36(1), 1–20.
- Sturm, T., Pumplun, L., Gerlach, J. P., Kowalczyk, M., & Buxmann, P. (2023). Machine learning advice in managerial decision-making: The overlooked role of decision makers' advice utilization. *The Journal of Strategic Information Systems*, 32(4), 101790. <https://doi.org/10.1016/j.jsis.2023.101790>
- Tai, K., Keem, S., Lee, K. Y., & Kim, E. (2024). Envy influences interpersonal dynamics and team performance: Roles of gender congruence and collective team identification. *Journal of Management*, 50(2), 556–587.
- Tangirala, S., & Ramanujam, R. (2008). Exploring nonlinearity in employee voice: The effects of personal control and organizational identification. *Academy of Management Journal*, 51(6), 1189–1203.
- Tilton, Z., LaVelle, J. M., Ford, T., & Montenegro, M. (2023). Artificial intelligence and the future of evaluation education: Possibilities and prototypes. *New Directions for Evaluation*, 2023(178–179), 97–109.
- Tong, S., Jia, N., Luo, X., & Fang, Z. (2021). The Janus face of artificial intelligence feedback: Deployment versus disclosure effects on employee performance. *Strategic Management Journal*, 42(9), 1600–1631.
- Tversky, A., & Kahneman, D. (1992). Advances in prospect theory: Cumulative representation of uncertainty. *Journal of Risk and Uncertainty*, 5(4), 297–323.
- Tversky, A., & Koehler, D. J. (1994). Support theory: A nonextensional representation of subjective probability. *Psychological Review*, 101(4), 547–567.
- Van den Heuvel, M., Demerouti, E., & Peeters, M. C. (2015). The job crafting intervention: Effects on job resources, self-efficacy, and affective well-being. *Journal of Occupational and Organizational Psychology*, 88(3), 511–532.
- Van Kleef, G. A. (2009). How emotions regulate social life: The emotions as social information (EASI) model. *Current Directions in Psychological Science*, 18(3), 184–188.
- Wang, Y. Y., & Chuang, Y. W. (2024). Artificial intelligence self-efficacy: Scale development and validation. *Education and Information Technologies*, 29(4), 4785–4808.
- Weiner, B. (1985). An attributional theory of achievement motivation and emotion. *Psychological Review*, 92(4), 548–573.
- Wesche, J. S., Hennig, F., Kollhed, C. S., Quade, J., Kluge, S., & Sonderegger, A. (2022). People's reactions to decisions by human vs. algorithmic decision-makers: The role of explanations and type of selection tests. *European Journal of Work and Organizational Psychology*, 33(2), 146–157.
- Wibberley, G., & Saundry, R. (2016). From representation gap to resolution gap: Exploring the role of employee voice in conflict management. In R., Saundry, P., Latreille, & I., Ashman (Eds.), *Selected papers on reframing resolution: Innovation and change in the management of workplace conflict* (pp. 127–148). London: Palgrave Macmillan UK.
- Wiesche, M., Pflügler, C., & Thatcher, J. B. (2024). The impact of social comparison on turnover among information technology professionals. *Journal of Management Information Systems*, 41(1), 297–324.
- Williams, S. H. (2020). AI advice: The irony of big data disclosures and the new advice paradigm. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.3657639>
- Wolsink, I., Den Hartog, D. N., Belschak, F. D., & Sligte, I. G. (2019). Dual cognitive pathways to voice quality: Frequent voicers improvise, infrequent voicers elaborate. *PLOS One*, 14(2), e0212608. <https://doi.org/10.1371/journal.pone.0212608>
- Wood, J. V. (1996). What is social comparison and how

- should we study it? *Personality and Social Psychology Bulletin*, 22(5), 520–537.
- Xavier, D. F., & Korunka, C. (2025). Integrating artificial intelligence across cultural orientations: A longitudinal examination of creative self-efficacy and employee autonomy. *Computers in Human Behavior Reports*, 18, 100623. <https://doi.org/10.1016/j.chbr.2025.100623>
- Yang, C., Bauer, K., Li, X., & Hinz, O. (2025). My advisor, her AI, and me: Evidence from a field experiment on human-AI collaboration and investment decisions. *Management Science*. Advance online publication. <https://doi.org/10.1287/mnsc.2022.03918>
- Yaniv, I. (2004). Receiving other people's advice: Influence and benefit. *Organizational Behavior and Human Decision Processes*, 93(1), 1–13.
- You, S., Yang, C. L., & Li, X. (2022). Algorithmic versus human advice: Does presenting prediction performance matter for algorithm appreciation?. *Journal of Management Information Systems*, 39(2), 336–365.
- Yue, B., & Li, H. (2023). The impact of human-AI collaboration types on consumer evaluation and usage intention: A perspective of responsibility attribution. *Frontiers in Psychology*, 14, 1277861. <https://doi.org/10.3389/fpsyg.2023.1277861>

The differential perception of leaders' response to employee-GenAI advice: A multi-level study based on social comparison view

HAN Yi¹, MA Zhaoyi², ZONG Shuwei³

(¹ School of Business Administration, Zhongnan University of Economics and Law, Wuhan 430073, China)

(² School of Economics and Management, Southwest Petroleum University, Chengdu 610500, China)

Abstract: The involvement of Generation artificial intelligence (GenAI) in organizational decision-making has become an irresistible trend. However, academic research on the adoption of GenAI advice remains insufficient. In the process of leadership decision-making, how do leaders compare and adopt employee advice, GenAI advice, and employee-GenAI team advice? What differences exist in leaders' perceptions? Grounded in social comparison theory, this study explores how leaders' differing perceptions shape their adoption of advice from employees, GenAI, and human-GenAI teams through five sub-studies: 1) How to define the GenAI advice-taking? 2) What are the differences in leaders' adoption of employee advice versus GenAI advice? 3) How do leaders differ in their adoption of employee-GenAI team advice? 4) How to compare the effect of employee-GenAI advice? 5) How do advice barriers in employee-GenAI comparisons affect advice quality, and to what extent can intervention strategies mitigate these effects? Ultimately, this study systematically constructs a multi-level model of leaders' adoption of employee-GenAI advice, providing a new perspective for interdisciplinary theoretical development and offering practical guidance for organizations to optimize human-AI collaborative decision-making and mitigate technological risks.

Keywords: leadership, GenAI advice, advice response, social comparison theory, perceptual differences