

《心理科学进展》审稿意见与作者回应

题目：AI 赋能视角下极限目标设置的驱动机制及其影响效应

作者：程佳琳，李劲松

第一轮

审稿人意见：

非常荣幸有机会审阅论文“AI 赋能视角下极限目标设置的驱动机制及其影响效应”。在这篇论文中，作者结合数智时代背景，探究了团队 AI 采用对极限目标设置的影响，以及 AI 在极限目标影响个体和团队创造力过程中的赋能作用。论文具有重要的探索价值，整体逻辑较为清晰。然而，在评审过程中，评审人也发现了一些需要关注的问题，包括理论、变量界定和逻辑推导等。本评审人将它们列在下面，希望对作者能有所帮助。

回应：非常感谢评审专家对我们文章提出的宝贵意见！您专业性的建议对提升该论文的质量具有非常大的帮助。在修改稿中，我们已依据您的评审意见，对问题一一做出了完善和修改，希望能得到您的肯定。

意见 1：理论逻辑与框架整合。选择“极限目标”的必要性和独特性。虽然论文区分了极限目标与其他目标，但目前的论证尚未有力地解释为什么一定是极限目标而非其他目标。作者需要更强有力地解释为什么 AI 的引入会导致领导者设置看起来不可能实现的极限目标，而不仅仅是其他目标，例如比较难的挑战性目标？AI 赋能视角下，极限目标对创造力的激发机制与其他目标（如挑战性目标）有何本质区别？

回应：非常感谢审稿专家提出的宝贵建议。针对您的建议，我们在“3.1 研究内容一：AI 赋能条件下极限目标的形成机理”部分，新增小节“3.1.1 团队 AI 采用对极限目标设置的影响”（正文第 19 页），详细解释 AI 的引入会促进领导者设置极限目标，而不是其他目标的原因。具体如下：

首先，AI 技术的引入会使领导意识到行业竞争已从“渐进式迭代”转向“颠覆性创新”，竞争对手可能会借助 AI 突破传统壁垒，形成“降维打击”（葛元骏，孙小蒙，2025; Alhur et al., 2025; Göhsl et al., 2025）。此时，若企业仅设置挑战性目标，难以应对行业竞争压力。而极限目标可以促进组织和成员跳出传统业务逻辑，重构技术路径和资源配置方式（Thompson et al., 1997; Kerr & LePelley, 2013），正是 AI 时代组织生存的核心需求。

其次，从目标设置条件来看，以往受到人力、经验等约束，成员的能力上限相对固定，领导者设置目标的期望是“在现有能力基础上适度拔高”，即挑战性目标（Locke & Latham, 1990）。但是 AI 引入后，领导意识到 AI 在数据处理、预测分析、信息搜寻等方面的强大潜

力 (Allal-Cherif et al., 2023; Bankins et al., 2024), 通过设置极限目标, 可实现“人力+AI”的能力倍增效应, 此时若仅设置挑战性目标, 则会限制 AI 赋能潜力的发挥。

最后, 极限目标相比于挑战性目标具有更高的不确定性和风险, 传统场景中领导者对极限目标的顾虑来源于过高的试错成本, 而挑战性目标风险可控是更为稳妥的选择 (Sitkin et al., 2011)。但是 AI 强大的实时监控能力可以在极限目标推进过程中纠正偏差 (Broekhuizen, 2023; Tang, 2023), 极大地降低极限目标的试错成本, 从而减少领导者对设置极限目标的风险顾虑。

本研究认为 AI 赋能视角下, 极限目标对创造力的激发机制与其他目标 (如挑战性目标) 的本质区别在于以下两点:

挑战性目标对创造力的激发机制是一种利用式学习。挑战性目标依据组织既往设定的目标、组织过往的绩效表现, 以及其他“可比参照”组织的历史绩效而设定 (Locke & Latham, 1990; Kerr & LePelley, 2013), 因此, 成员尽管需要付出额外努力, 但无需突破现有路径只需对已有知识、技术、方法进行精炼与优化 (Huang & Li, 2012; Neill et al., 2020)。在该过程中, 挑战性目标未超越成员的能力上限 (Locke & Latham, 2002; Vainieri et al., 2016), AI 仅作为辅助工具提升成员的工作效率和缓解工作压力, 人机关系是“主从关系”, 成员仍然是该机制中的核心主体。

极限目标对创造力的激发机制是探索性学习, 极限目标因其“看似不可能性”对成员构成强烈的愿景吸引 (Sitkin et al., 2011; Kerr & LePelley, 2013), 激发成员内在动机和对新领域、新方法、新知识的探索与尝试, 从而促进成员产生创造力。同时, AI 赋能也能给成员提供强大的“能力支撑”, 在该过程中, AI 不仅是工具, 更是成员的“工作伙伴”, 让成员避免了因目标过高而产生挫败感 (O'Neill et al., 2023; Tsai et al., 2022)。这种人机合作是极限目标影响创造力的独有机制, 挑战性目标因无需突破成员的能力上限, 推进过程仍以成员为主导, 则无法激活人机合作的价值。

为呈现该部分内容, 本研究增设小节“2.2 极限目标与挑战性目标对创造力激发机制的区别。具体请见正文 15-16 页。

参考文献

- 葛元骏, 孙小蒙. (2025). 数智化转型对科技型企业高端颠覆性创新的影响——基于新一代人工智能创新发展试验区试点的准自然实验[J/OL]. 工程管理科技前沿. <https://link.cnki.net/urlid/34.1336.N.20251231.1055.002>
- Alhur, M., Omar, F., Alqirem, R., & Yousuf, A. A. N. (2025). Adopting the future: How generative AI, agentic AI, and blockchain are redefining success for dubai's emerging start-ups. *Human Behavior and Emerging Technologies*, 20, 9751686.
- Allal-Ch'erif, O., Climent, J. C., & Berenguer, K. J. U. (2023). Born to be sustainable: How to combine strategic disruption, open innovation, and process digitization to create a sustainable business. *Journal of Business Research*, 154, 113379.
- Bankins, S., Ocampo, A. C., Marrone, M., Restubog, S. D., & Woo, S. E. (2024). A multilevel review of artificial intelligence in organizations: Implications for organizational behavior research and practice. *Journal of*

- Organizational Behavior*, 45(24), 159–182.
- Broekhuizen, T., Dekker, H., de Faria, P., Firk, S., Nguyen, D. K., & Sofka, W. (2023). AI for managing open innovation: Opportunities, challenges, and a research agenda. *Journal of Business Research*, 167, 114196.
- Göhlsl, J. F. (2025). What if disruption really happens—Are competition law and digital regulation fit for a new era of AI-driven competition? *Journal of Competition Law & Economics*, 1–29.
- Kerr, S., & Lepak, D. (2013). Stretch goals: Risks, possibilities, and best practices. In E. Locke, & G. P. Latham (Eds.). *New developments in goal setting and task performance* (pp. 21 – 33). New York: Routledge.
- Locke, E. A., & Latham, G. P. (1990). A theory of goal setting and task performance. Englewood Cliffs, NJ: Prentice Hall.
- Locke, E. A., & Latham, G. P. (2002). Building a practically useful theory of goal setting and task motivation: A 35-year odyssey. *American Psychologist*, 57(9), 705–717.
- O’Neill, T. A., Flathmann, C., McNeese, N. J., & Salas, E. (2023). Human-autonomy teaming: Need for a guiding team-based framework?. *Computers in Human Behavior*, 146, 107762.
- Sitkin, S. B., See, K. E., Miller, C. C., Lawless, M. W., & Carton, A. M. (2011). The paradox of stretch goals: Organizations in pursuit of the seemingly impossible. *Academy of Management Review*, 36(3), 544–566.
- Tang, P. M., Koopman, J., Yam, K. C., Cremer, D. D., Zhang, J. H. & Reynders, P. The self-regulatory consequences of dependence on intelligent machines at work: Evidence from field and experimental studies. *Human Resource Management*, 62, 721–744.
- Thompson, K. R., Hochwarter, W. A., & Mathys, N. J. (1997). Stretch targets: What makes them effective? *Academy of Management Executive*, 11(3), 48–60.
- Tsai, C. Y., Marshall, J. D., Choudhury, A., Serban, A., Hou, Y.Y.T., Jung, M. F., Dionne, S. D., & Yammarino, F. J. (2022). Human-robot collaboration: A multilevel and integrated leadership framework. *The Leadership Quarterly*, 33, 101594.
- Vainieri, M., Volaa, F., Soriano, G. G., & Nuti, S. (2016). How to set challenging goals and conduct fair evaluation in regional public health systems. Insights from Valencia and Tuscany Regions. *Health Policy*, 120, 1270–1278.

意见 2：理论逻辑与核心构念存在矛盾。论文沿用了过去的极限目标定义，即“依据成员当前的实践、技能和知识似乎无法实现的目标”。同时，论文认为是 AI 作为一种强大的资源“打破了极限目标实现的约束条件”。这里存在矛盾：如果 AI 已经赋予了成员完成任务的能力和资源，那么该目标对于拥有 AI 的团队而言，是否还符合“似乎无法实现”的定义？作者需要更深刻地阐述：在 AI 赋能条件下，极限目标的内涵是否发生了变化？是否是因为 AI 的引入，导致管理者设定了阈值极高的新目标，从而维持了其“极限”的属性？

回应：非常感谢评审专家对理论逻辑与核心构念的质疑，您的疑问为本文厘清理论逻辑与优化构念提供了关键指导。针对您提出的问题，首先向您说明本文的表述初衷：并不是阐述 AI 已赋予成员完成任务的能力和资源，而是强调 AI 在成员推进极限目标过程中的动态赋能作用。此次矛盾的产生源于本文对 AI 赋能的阶段没有定义清楚以及不当的表达，导致核心逻辑未能清晰传递。

本文所指的 AI 赋能是在成员推进极限目标的过程中，AI 给予成员的动态性的能力和专业支持，而非在目标设定之初，就为成员提供了完成该目标的完整能力和资源。在极限目标设定阶段，极限目标设定的基准仍锚定成员“当前的实践、技能和知识”，此时 AI 尚未参与

目标推进过程。在极限目标推进阶段，AI 才作为一种强大的资源介入，针对成员的能力和知识短板进行动态赋能（张志学等, 2024; Anthony et al., 2023; Bankins et al., 2024; O'Neill et al., 2023），在该过程中也只是增强了团队成员推进目标的能力，而不是直接将目标的“不可实现”变成“轻易可实现”。

因此，本文希望极限目标的内涵仍沿用经典定义，并对 AI 赋能的发生阶段进行解释且删除“打破了极限目标实现的约束条件”等不当表述，从而清晰传递核心逻辑。**该观点在论文的“1.问题提出”部分呈现，具体可见论文 13 页。**

参考文献

- 张志学, 华中生, 谢小云. (2024). 数智时代人机协同的研究现状与未来方向. *管理工程学报*, 38(1), 1–13.
- Anthony, C., Bechky, B. A., Fayard, A. L. (2023) “Collaborating” with AI: Taking a system view to explore the future of work. *Organization Science*, 34(5), 1672–1694.
- Bankins, S., Ocampo, A. C., Marrone, M., Restubog, S. D., & Woo, S. E. (2024). A multilevel review of artificial intelligence in organizations: Implications for organizational behavior research and practice. *Journal of Organizational Behavior*, 45(24), 159–182.
- O'Neill, T. A., Flathmann, C., McNeese, N. J., & Salas, E. (2023). Human-autonomy Teaming: Need for a guiding team-based framework?. *Computers in Human Behavior*, 146, 107762.

意见 3：理论框架的整合性不足。作者在研究内容二中，使用了目标设置理论解释成员探索性学习在极限目标和成员创造力的中介作用，使用了社会-技术系统理论解释 AI 有用性感知在其中的调节作用。这种做法显得理论框架较为碎片化。未能在一个统一的理论下解释所有变量关系，削弱了理论贡献的深度。因此，建议作者尝试用一个统一的理论解释整个模型，增强各变量之间逻辑关系的连贯性。

回应：感谢评审专家的专业性建议。在您的提点下，我们再次回顾社会-技术系统理论，发现该理论可以完整解释研究二的理论模型，具体理由如下：

社会-技术系统理论将团队视为由两个相互独立但又彼此关联的系统构成：一是由身处各类关系中的人所组成的社会系统，二是包含完成任务所需全部工具与流程的技术系统（Cummings et al., 1978; Bostrom & Heinen, 1977）。高强度的任务环境刺激，会促进社会系统和技术系统相互整合，从而实现团队的最优绩效（Cummings et al., 1978; Leung & Wang, 2015）。极限目标作为高强度的任务环境刺激，会推动团队社会系统和技术系统的相互整合优化，即成员主动进行探索性学习，尝试新的工作工具、协作流程和分析方法（Gupta et al., 2006; Huang & Li, 2012; Neill et al., 2020），该过程会促进兼具新颖性和可行性的创造性方案的产生，即创造力（Chang & Shih, 2019; Escriba-Carda et al., 2017）。

此外，随着社会-技术系统理论的发展，技术系统由原先的工具与流程，扩充到各种不同类型的科技等更为开放的系统（Abadie et al., 2023; Leung & Wang, 2015; Makarius et al., 2020）。同时，该理论强调在工作系统中，技术对个体实现系统目标的协助作用（Appelbaum et al., 1997; Makarius et al., 2020）。虽然极限目标触发了成员“想探索”的意愿，但技术系统影

响成员“能探索”的程度。可用性较高的技术可为成员进行探索性学习创造有利条件，让极限目标触发的探索性学习行为更顺利地开展，进而推动创造力提升。

因此，本研究将基于社会-技术系统理论对研究二的理论模型进行论述。结合您提出的意见 3，修改后的内容请见正文“3.2 研究内容二：极限目标对个体创造力的影响机制研究”部分（第 22-24 页）。

参考文献

- Abadie, A., Roux, M., Chowdhury, S., & Dey, P. (2023). Interlinking organizational resources, AI adoption and omnichannel integration quality in Ghana's healthcare supply chain. *Journal of Business Research*, 162, 113866.
- Appelbaum, S. H. (1997). Socio-technical systems theory: an intervention strategy for organizational development, *Management Decision*, 35 (6), 452–463.
- Bostrom, R. P., & Heinen, J. S. (1977). MIS problems and failures: a socio-technical perspective. *MIS Quarterly*, 1(4), 11–28.
- Chang, Y. Y., & Shih, H. Y. (2019). Work curiosity: A new lens for understanding employee creativity. *Human Resource Management Review*, 29, 100672.
- Cummings, T. G. (1978). Self-regulating work groups: A socio-technical synthesis. *Academy of Management Review*, 3, 625–634.
- Escriba-Carda, N., Balbastre-Benavent, F., & Canet-Giner, M. T. (2017). Employees' perceptions of high-performance work systems and innovative behavior: The role of exploratory learning. *European Management Journal*, 35, 273–281.
- Gupta, A. K., Smith, K. G., & Shalley, C. E. (2006). The interplay between exploration and exploitation. *Academy of Management Journal*, 49(4), 693–706.
- Huang, J. W., & Li, Y. H. (2012) Slack resources in team learning and project performance. *Journal of Business Research*, 65, 381–388.
- Leung, K. & Wang, J. (2015). Social processes and team creativity in multicultural teams: A socio-technical framework. *Journal of Organizational Behavior*, 36, 1008–1025.
- Makarius, E. E., Mukherjee, D., Fox, J. D. & Fox, A. K. (2020) Rising with the machines: A sociotechnical framework for bringing artificial intelligence into the organization. *Journal of Business Research* 120, 262–273.
- Neill, S., Pathak, R. D., Ribbens, B. A., Noel, T. W., & Singh G. (2020). The influence of managerial optimism and self-regulation on learning and business growth expectations within an emerging economic context. *Asia Pacific Journal of Management*, 37(1), 187–204.

意见4：变量界定的准确性。关于AI智能化水平。作者提出，“AI的智能化水平越高，团队AI采用和团队领导创新期待之间的正向关系越强”。然而，根据文章中的论述逻辑，影响领导反应的并非是AI客观的技术指标（智能化水平），而是团队领导对AI能力的主观认知，建议将变量名称修正为“AI智能化水平感知”。

回应：感谢评审专家提出的宝贵意见。根据您的建议，我们已将文中的“AI 智能化水平”都修正为“AI 智能化水平感知”。

意见 5：作用机制中的竞争性假设与负面效应。“极限目标”对学习行为可能存在负向影响。

作者在研究内容二中，基于目标设置理论，提出“极限目标会促进成员实施探索性学习”。然而，目标设置理论同样指出，当目标难度超过一定阈值，个体可能会降低自我效能感，导致放弃努力或降低动机，而非积极探索。作者指出，“极限目标对知识和技能的范围要求已经远超成员目前所掌握的”。因此，极限目标也可能阻碍成员实施探索性学习。作者需要考虑被忽视了的目標過難可能帶來的阻礙效應。

回应：感谢评审专家的启发性提问。围绕该问题，我们再一次去研读了相关的理论文献，发现目标设置理论的提出者 Locke 和 Latham 在 2013 年出版的《New Developments in Goal Setting and Task Performance》这本书，为研究二的理论模型提供了强有力的解释。

根据目标设置理论，当目标难度超过一定阈值，个体会降低自我效能感，导致放弃努力或降低动机，这一观点提示领导者需警惕给员工无意间设定过高的工作目标（Locke & Latham, 1984）。而极限目标则跳出了目标设置理论的框架，因为极限目标并非是管理者“必须警惕防范”，无意间设定的过高目标，而是被刻意设定在这样的水平（Kerr & Lepelley, 2013; Sitkin et al., 2011），以促进员工进行更多的探索性行为，催生员工提出更多非传统的解决方案。极限目标带有组织的“愿景属性”，能够将人们的注意力转向未来的各种可能，甚至激发更强劲的行动动力（Kerr & Lepelley, 2013; Sitkin et al., 2011）。

此外，大量实践案例表明极限目标能通过激励员工实施实验探索、创新实践等方式产生更多创新性的成果，这一目标形式也被通用电气、联邦健康集团、丰田等诸多管理成熟、秉持以人为本理念的优秀企业所采用（Ahmadi et al., 2022; Kerr & Lepelley, 2013; The Economist, 2011）。同时，领导者设置极限目标的初衷是引导员工以创新性的方式重新认知自身工作并开展各项任务。在实践中，极限目标并非对常规目标的替代，而是作为补充手段存在，且员工不会因尝试极限目标的失败而受到惩罚，这也从实施层面避免了因知识技能暂未匹配而产生的自我效能感降低、动机退缩等问题（Kerr & Landauer, 2004; Kerr & Lepelley, 2013; Sitkin et al., 2011）。综上所述，极限目标承载着组织的愿景期许，能够有效激发员工的内在工作动机和工作干劲，进而推动其开展探索性学习，实现创造力提升。此外，极限目标会导致员工过于投入工作而负向影响身心健康和工作家庭平衡（Kerr & Lepelley, 2013）。

结合意见 1 的第三点建议，本文将该部分内容纳入“3.2.1 极限目标对员工探索性学习的正向影响”部分，请见正文 22-23 页。

参考文献

- Ahmadi, S., Jansen, J. J. P., & Eggers, J. P. (2022). Using stretch goals for idea generation among employees: One size does not fit all. *Organization Science*, 33(2), 671–687.
- Kerr, S., & Landauer, S. (2004). Using stretch goals to promote organizational effectiveness and personal growth: General electric and Goldman Sachs. *Academy of Management Executive*, 18(4), 134–138.
- Kerr, S., & Lepelley D. (2013). Stretch goals: Risks, possibilities, and best practices. In E. Locke, & G. P. Latham (Eds.). *New developments in goal setting and task performance* (pp. 21–33). New York: Routledge.
- Locke, E. A., & Latham, G. P. (1984). *Goal setting: A motivational technique that works*. Englewood Cliffs, NJ: Prentice Hall.
- Sitkin, S. B., See, K. E., Miller, C. C., Lawless, M. W., & Carton, A. M. (2011). The paradox of stretch goals:

Organizations in pursuit of the seemingly impossible. *Academy of Management Review*, 36(3), 544–566.
The Economist. A work-out for corporate America. Academic OneFile. Web. June 28, 2011.

意见 6: “AI 有用性”可能带来技术依赖风险。作者在研究内容二提出, “AI 有用性感知程度越高, 极限目标和成员探索性学习之间的正向关系越强”。这一命题忽视了人性的“认知吝啬”倾向。读者可能会疑惑: AI 有用性感知高, 是否一定促进探索性学习? 可能会出现一种情况: 如果 AI 有用性很高, 成员可能会产生技术依赖或认知懒惰, 直接照搬 AI 生成的方案(利用性行为), 从而通过减少付出完成任务, 反而减少了深度的探索性学习。作者需要在理论中排除这一竞争性假设。

回应: 非常感谢您提出的具有批判价值的问题, 经深入思考, 我们发现本研究在概念界定与假设逻辑的阐释上存在疏漏。本研究将基于以下观点, 排除竞争性假设:

极限目标的核心特征是目前没有已知的达成路径和标准化的解决方案 (Ahmadi et al., 2022; Sitkin et al., 2011), 其任务属性决定 AI 无法生成直接可以照搬的现成方案。此外, 从 AI 的能力角度来说, AI 在产品研发、技术创新等高阶领域还是存在非常大的短板, 无法超越人类强大的大脑 (Benedikt et al., 2023; Haefner et al., 2021)。由此, 本研究提到的 AI 赋能体现为 AI 对其他成员的能力和知识短板进行动态赋能, 在该过程中增强其他成员推进目标的能力, 而不是直接将目标的“不可实现”变成“轻易可实现”。在此情境下, AI 有用性则指 AI 可为成员提供大数据分析、知识线索等的人机合作性价值, 而非提供终极答案 (Glikson & Woolley, 2020)。因此, 成员对 AI 有用性的高感知, 将会强化其借助 AI 开展探索性学习的动机, 而非替代探索性学习。

本研究将该论证纳入正文的“3.2.3AI 有用性感知的调节作用”的部分, 具体修改请见正文第 23-24 页。

参考文献

- Ahmadi, S., Jansen, J. J. P., & Eggers, J. P. (2022). Using stretch goals for idea generation among employees: One size does not fit all. *Organization Science*, 33(2), 671–687.
- Benedikt, M., Daniel, R., & Matthias, K. (2023). Barriers to the use of artificial intelligence in product development—A survey of dimensions involved. *Proceedings of the Design Society*, 3, 757–766.
- Glikson, E., & Woolley, A. W. (2020). Human trust in artificial intelligence: Review of empirical research. *Academy of Management Annals*, 14, 627–660.
- Haefner, N., Wincent, J., Parida, V., & Gassmann, O. (2021). Artificial intelligence and innovation management: A review, framework, and research agenda. *Technological Forecasting & Social Change* 162, 120392.
- Sitkin, S. B., See, K. E., Miller, C. C., Lawless, M. W., & Carton, A. M. (2011). The paradox of stretch goals: Organizations in pursuit of the seemingly impossible. *Academy of Management Review*, 36(3), 544–566.

意见 7: “人-AI 合作”中可能存在防御性心理。作者在研究内容三提出, “极限目标对人-AI 合作具有显著的正向影响”。然而, 在极限目标的高压下, 成员有可能会产生防御性心理, 将 AI 视为推卸责任的对象, 或者因算法不透明而产生不信任, 从而阻碍与 AI 的合作。因

此，作者需要解释极限目标的高压是否会阻碍人与 AI 的有效磨合。

回应：非常感谢评审专家的专业意见！我们非常认同，在无情境约束的条件下，极限目标的确可能诱发防御性心理，进而阻碍人-AI 合作。研究三中，人-AI 合作被界定为在人机混合团队中，人和 AI 各自承担相应的工作角色，并分享对任务的理解，讨论交流不同观点和共同实现团队目标的过程（O'Neill et al., 2023; Tsai et al., 2022）。本研究认为人-AI 工作角色清晰是消解极限目标高压下防御性心理、保障人-AI 合作的关键边界条件：当团队内人和 AI 的角色定位、任务分工与工作权责被清晰界定后，不仅从根源上消除成员将 AI 视为责任推诿对象的可能，也有效降低因算法认知模糊引发的不信任感。已有研究亦证实，高水平的人-AI 角色清晰度能推动成员精准感知人机双方的职责边界，进而建立对 AI 的合作信任（Chowdhury et al., 2022; Frey & Osborne, 2017）。因此，在高角色清晰度的情境下，团队成员面对极限目标时，更易形成对 AI 的正向认知，进而主动开展人-AI 合作行为。

基于此，研究三的假设 1 直接提出二者的交互效应：人-AI 角色清晰正向调节极限目标和人-AI 合作的关系，即人-AI 角色清晰水平越高，极限目标与人-AI 合作之间的正向关系越强。具体的修改请见正文 25-26 页。

参考文献

- 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.
- Frey, C. B., & Osborne, M. A. (2017). The future of employment: How susceptible are jobs to computerisation? *Technological Forecasting and Social Change*, 114, 254–280.
- Tsai, C. Y., Marshall, J. D., Choudhury, A., Serban, A., Hou, Y.Y.T., Jung, M. F., Dionne, S. D., & Yammarino, F. J. (2022). Human-robot collaboration: A multilevel and integrated leadership framework. *The Leadership Quarterly*, 33, 101594.
- O'Neill, T. A., Flathmann, C., McNeese, N. J., & Salas, E. (2023). Human-autonomy Teaming: Need for a guiding team-based framework?. *Computers in Human Behavior*, 146, 107762.

意见 8：研究内容四的定位。作者计划在研究内容四中进行极限目标形成机理及其影响效应的嵌入式案例研究，目的是“探索数智时代极限目标形成机理及其影响效应”。但是，考虑到研究一至三中已经构建了较为详尽的理论模型，那么研究四是否不再归类为探索性研究会比较合适？如果它是为了探索新机理，那么它应该置于研究的最开始。建议作者明确界定该案例研究的独特价值（例如，它是否能揭示量化数据无法捕捉的过程黑箱）。

回应：感谢评审专家的宝贵意见。基于您的建议，我们对该问题进行了思考，非常认同您的观点：研究一至研究三已构建较为系统完善的理论模型，研究四无需再开展探索性案例研究。我们最初的研究设想，是想了解所提出的理论模型在实际工作场景中的可行性，经相关文献查阅后发现（刘楠等, 2025; 裴冠雄等, 2026），田野实验方法能够更精准地实现这一研究意图，推动理论成果向实地应用落地。因此，研究四拟从实际应用层面，将前述研究中“AI 采用→团队极限目标设置→个体和团队创造力”的主要结论在实际场景中进行运用，拟通过

田野实验评估实际效果。具体内容如下：

基于上述三个研究，本项目已构建较为详尽的理论模型。研究四拟从实际应用层面，将前述研究中“环境刺激→目标设置→成员行为结果”的主要结论在实际场景中进行运用，拟通过田野实验进一步验证数智时代极限目标形成机理及其影响效应，为企业促进个体和团队创造力提供参考。

研究四拟选取一家处于数智化转型阶段且具备团队协作属性的企业作为合作对象进行田野实验。首先，对团队 AI 采用进行操纵：将该企业多个创新团队分为干预组（AI 采用组）：团队可使用指定的 AI 工具，并接受基础的 AI 使用培训；对照组（无 AI 采用组）：团队保持常规工作模式，不使用指定 AI 工具。其次，对 AI 采用团队和无 AI 采用团队的目标设置流程、规则、目标类型、时间节点做统一介绍，但完全不干预目标的难度和新颖性，让团队在各自的工作条件下自主决策。最后，采用问卷、访谈等方式收集理论模型中观测变量的相关数据。田野实验既能借助真实情景验证研究结果的稳定性，进一步强化理论的外部效度；又能推动理论成果向实际应用落地（Harrison & List, 2004; 刘楠等, 2025; 裴冠雄等, 2026）。

具体的修改请见论文 27 页。

参考文献

- 刘楠, 邱钰婷, 董志强, 李爱梅. (2025) 抚育动机何以促进亲社会行为? 基于进化适应性的认知-能量机制及助推. *心理科学进展*, 33(11), 1854–1869.
- 裴冠雄, 董波, 金佳, 孟亮, 张加林. (2026). 营销数字人对话智能特征的动态加工与神经机制. *心理科学进展*, 34(2), 227–238.
- Harrison, G. W., & List, J. A. (2004). Field experiments. *Journal of Economic Literature*, 42(4), 1009–1055.

意见 9：写作与规范。作者在论文中“人工智能”与“AI”混用。建议在正文首次出现“人工智能（Artificial Intelligence, AI）”后，后文统一使用“AI”，以保持学术语言的简洁与规范。

回应：感谢评审专家提出的修改意见。根据您的建议，我们已解决论文中“人工智能”与“AI”混用的问题。

第二轮

审稿人意见：

作者对审稿意见的回应态度认真、修改充分、论证扎实，尤其在理论框架统一、概念矛盾澄清、竞争性假设排除等方面，显示出作者对研究问题深入的理解与严谨的学术态度。修改后的论文无论是在理论逻辑、变量界定，还是研究设计上均有显著提升，初步达到期刊发表要求，建议修改后录用。当然，尽管整体修改质量较高，仍有以下几点可供作者进一步思考。

回应：非常感谢专家对本文修改工作的细致审阅与肯定。针对专家提出仍可进一步思考与优化的内容，我们已认真研读与分析，并对相关问题逐一进行了完善与修改，力求进一步提升论文质量，恳请专家予以审阅指正。

意见 1：田野实验的具体设计。研究四虽调整了定位，但对实验的样本、干预方式、观测变量、时间跨度等细节尚未充分展开。若能在正文中补充初步的实验设计思路，将增强研究四的可操作性与说服力。

回应：非常感谢专家对研究四提出的建议。根据您的建议，我们详细阐述了田野实验计划选取的企业样本、干预方式、观测变量、时间跨度等细节，具体内容如下：

基于上述三个研究，本文已构建较为详尽的理论模型。研究四拟从实际应用层面，将前述研究中“环境刺激→目标设置→员工行为结果”的主要结论在实际场景中进行运用，拟通过田野实验进一步验证数智时代极限目标形成机理及其影响效应，为企业促进个体和团队创造力提供参考。具体实验设计如下：

本研究计划寻找一家科技公司作为合作对象，该公司应具备产品设计、运营、营销策划等创意性团队。选择 30 个 4-6 人的团队参与该田野实验，随机将团队分配至干预组（AI 采用组）或对照组（无 AI 采用组）。在干预组中，为团队配备飞书软件，并将“OpenClaw 龙虾”接入飞书，确保干预组每个团队在飞书工作群中添加 OpenClaw。并对团队成员开展为期两周的标准化 OpenClaw 操作培训，确保成员掌握基本操作与功能边界。对照组中的团队保持原有工作模式，不使用上述 AI 工具，也不接受相关培训。

本田野实验总周期为 4 个月，预期分三个阶段。时间点 1（干预前 1 个月）：完成样本筛选，收集团队当前目标水平、创造力水平、团队规模、团队存续时间、团队领导任期等数据。时间点 2（共两个月）：随机将团队分配至干预组（AI 采用组）或对照组（无 AI 采用组）。干预组完成 AI 工具部署与培训，进入正常任务周期。两组团队在各自条件下开展实际工作，研究者不干预目标设置与任务执行过程，仅跟踪记录 AI 使用情况与团队目标设置情况，并每隔两周让员工和领导填写团队 AI 采用、AI 智能化水平、领导创新期待和极限目标设置等量表，并结合访谈了解团队目标推进过程中的关键事件与机制表现。时间点 3（1 个月之后）：统一收集个体创造力、团队创造力、个体探索性学习、人-AI 合作、AI 有用性感知、人-AI 工作角色清晰等变量的数据，并收集与创造力相关的客观绩效数据。

关于变量的测量。团队 AI 采用和 AI 智能化水平分别依据 Cheng et al. (2023)和 Balakrishnan et al. (2022)开发的量表进行改编；团队领导创新期待和极限目标设置分别参考 Carmeli 和 Schaubroeck (2007)以及 Zhang 和 Jia (2013)开发的量表进行测量。AI 有用性感知和人-AI 合作参考 Fernandes 和 Oliveira (2021)和 Kong et al. (2023)开发的量表。人-AI 工作角色清晰的测量依据 Chowdhury 等 (2022)开发的成熟量表。探索性学习、个体创造力和团队创造力分别参考 Neill et al.(2020)、Farmer et al. (2003)和 Jia et al.(2014)采用的量表。

田野实验既能借助真实情景验证研究结果的稳定性，进一步强化理论的外部效度；又能推动理论成果向实际应用落地（Harrison & List, 2004; 刘楠等, 2025; 裴冠雄等, 2026）。基于此，研究四提出以下命题：

命题 13：在真实工作场景中，团队 AI 采用会推动团队极限目标设置并促进个体和团队创造力。

该内容呈现在正文“3.5 研究内容四：极限目标形成机理及其影响效应的应用验证”部分（正文第 31-32 页）。

参考文献

- 刘楠, 邱钰婷, 董志强, 李爱梅. (2025) 抚育动机何以促进亲社会行为? 基于进化适应性的认知-能量机制及助推. *心理科学进展*, 33(11), 1854–1869.
- 裴冠雄, 董波, 金佳, 孟亮, 张加林. (2026). 营销数字人对话智能特征的动态加工与神经机制. *心理科学进展*, 34(2), 227–238.
- Balakrishnan, J., Abed, S. S., & Jones, P. (2022). The role of meta-UTAUT factors, perceived anthropomorphism, perceived intelligence, and social self-efficacy in chatbot-based services?. *Technological Forecasting & Social Change*, 180, 121692.
- Carmeli, A., & Schaubroeck, J. (2007). The influence of leaders' and other referents' normative expectations on individual involvement in creative work. *The Leadership Quarterly*, 18(1), 35–48.
- Cheng, B., Lin, H. X., & Kong, Y. R. (2023). Challenge or hindrance? How and when organizational artificial intelligence adoption influences employee job crafting. *Journal of Business Research*, 164, 113987.
- 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.
- Farmer, S. M., Tierney, P., & Kung-McIntyre, K. (2003). Employee creativity in Taiwan: An application of role identity theory. *Academy of Management Journal*, 46, 618 – 630.
- Fernandes, T., & Oliveira, E. (2021). Understanding consumers' acceptance of automated technologies in service encounters: Drivers of digital voice assistants adoption. *Journal of Business Research*, 122, 180–191.
- Harrison, G. W., & List, J. A. (2004). Field experiments. *Journal of Economic Literature*, 42(4), 1009–1055.
- Jia, L., Shaw, J. D., Tsui, A. S., & Park, T. Y. A. (2014). Social-structural perspective on employee-organization relationships and team creativity. *Academy of Management Journal*, 57(3), 869–891.
- Kong, H. Y., Yin, Z. H., Baruch, Y. H., & Yuan, Y. (2023). The impact of trust in AI on career sustainability: The role of employee-AI collaboration and protean career orientation. *Journal of Vocational Behavior*, 146, 103928.
- Neill, S., Pathak, R. D., Ribbens, B. A., Noel, T. W., & Singh G. (2020). The influence of managerial optimism and self-regulation on learning and business growth expectations within an emerging economic context. *Asia*

Pacific Journal of Management, 37(1), 187–204.

Zhang, Z., & Jia, M., (2013). How can companies decrease the disruptive effects of stretch goals? The moderating role of interpersonal- and informational-justice climates. *Human relations*, 66(7), 993–1020.

意见 2：人-AI 角色清晰的测量与操纵。该变量在理论中承担关键的边界作用，但目前在文中尚未明确其在研究三中拟采用的测量工具或操纵方式。建议后续补充具体操作化方案。

回应：非常感谢评审专家提出的宝贵意见。本研究对人-AI 角色清晰性的测量拟参考 Chowdhury 等（2022）开发的成熟量表，该量表共包含 10 个题项，具体内容可详见该文献。考虑到全文结构的整体性与连贯性，并结合您提出的第一条建议，本研究已将人-AI 角色清晰性的测量方式与其他变量统一整合在研究四中进行呈现。

参考文献

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.

意见 3：跨层次理论的进一步衔接。研究一聚焦团队层面（团队 AI 采用→领导认知→极限目标），研究二、三分别进入个体与团队层面。虽然理论框架统一为社会-技术系统理论，但跨层次逻辑的衔接仍可进一步强化，例如补充“领导设置的极限目标如何被团队与个体感知并转化为行为”的机制描述。

回应：非常感谢专家提出的建设性意见。根据您的建议，我们补充了“领导设置的极限目标如何被团队与个体感知并转化为行为”的机制描述，具体内容如下：

3.1 团队极限目标被个体与团队感知并转化为行为的机制描述

研究一聚焦于团队极限目标设置的形成机理，研究二和研究三分别探讨团队极限目标对个体创造力和团队创造力的影响，本文将进一步明确三者之间的跨层次逻辑。具体而言，团队领导在设置团队极限目标后，通常会通过正式的任务分配、团队会议、绩效沟通等方式，将该目标的内容、要求等传递给团队成员（Vanderstukken et al., 2019; Zhou et al., 2019）。在这一过程中，成员基于自身对目标的理解和接收到的信息，形成个体层面的极限目标感知。它作为任务情境刺激，激发员工开展探索性学习，进而提升个体创造力。这一路径反映团队目标设置影响到个体认知与行为向下传递的过程。

同时，团队成员之间的互动与意义构建，会促使团队层面涌现出对极限目标的共享认知（Chan, 1998; Naumann & Bennett, 2000; Salancik & Pfeffer, 1978），即团队成员对极限目标难度与新颖性的集体理解，表现为团队层面的目标共识。它作为团队层面的任务刺激，通过促进人-AI 合作这一过程，提升团队创造力。这一路径体现团队层面的极限目标共识对团队整体协作模式与结果的影响。

该内容添加在正文“3.1 团队极限目标被个体与团队感知并转化为行为的机制描述”部分

（正文第 23 页）。

参考文献

- Chan, D. (1998). Functional relations among constructs in the same content domain at different levels of analysis: A typology of composition models. *Journal of Applied Psychology*, 83(2), 234–246.
- Naumann, S. E., & Bennett, N. (2000). A case for procedural justice climate: Development and test of a multilevel model. *Academy of Management Journal*, 43(5), 881–889.
- Salancik, G. R., & Pfeffer, J. (1978). A social information processing approach to job attitudes and task design. *Administrative Science Quarterly*, 23(2), 224–253.
- Vanderstucken, A., Schreurs, B., Germeys, F., Van den Broeck, A., & Proost, K. (2019). Should supervisors communicate goals or visions? The moderating role of subordinates' psychological distance. *Journal of Applied Social Psychology*, 49(11), 671–683.
- Zhou, L., Wang, M., & Vancouver, J. B. (2019). A formal model of leadership goal striving: Development of core process mechanisms and extensions to action team context. *Journal of Applied Psychology*, 104(3), 388–410.

编委意见：经过两轮的修改，文章基本已经达到发表水平。建议前面文献综述做适当删减，这样内容更加精炼。