

和而不同：生成式人工智能凸显下 人类的社会创造策略*

周 详 白博仁 张婧婧 刘善柔

(南开大学社会学院社会心理学系, 南开大学前沿交叉学科研究院, 天津 300350)

摘 要 横空出世的生成式人工智能(AI)模糊了 AI 与人类心智间的界限, 人们会如何化解这种相似性值得关注。基于社会认同理论, 通过 4 项递进研究, 系统探讨了生成式 AI 凸显下个体社会创造策略的使用情况, 并考察了人类中心主义的调节作用。结果发现: 个体会创造性地定义人类心智与生成式 AI 的区分方式(研究 1), 并运用此种社会创造策略来应对生成式 AI 凸显(研究 2~4)。此外, 人类中心主义会增强生成式 AI 凸显情境下个体的社会创造策略使用程度(研究 3~4)。研究结果揭示了生成式 AI 凸显情境下个体应用社会创造策略的具体表现形式及其积极意义, 阐明了人类中心主义潜在的建设性作用, 并为智能时代如何借助社会创造策略巧妙推动人类与 AI 和谐共生提供了新的启发。

关键词 生成式人工智能, 人类心智, 社会创造策略, 人类中心主义

分类号 B849: C91

在人工智能时代, 唯有更深入地理解人类的独特本质, 才能掌握我们的未来。

——莱斯利·瓦利安特(Leslie Valiant, 图灵奖得主)

1 引言

人工智能(Artificial Intelligence, AI)技术的发展始终伴随着人们对于难以区分 AI 与人类的担忧(Cha et al., 2020; Ferrari et al., 2016)。而生成式 AI 技术的出现让 AI 与人类变得更加难以分辨(Mitchell, 2024), 进一步加剧了这种忧虑(Millet et al., 2023)。与传统仅能进行数据分类、预测等机械性任务的判别式 AI 不同(Banh & Strobel, 2023), 生成式 AI 在连续对话、内容生成、整合记忆等方面的出色表现不仅极大地动摇了人们关于人类独特性的既有认知(Santoro & Monin., 2023), 同时也更加助长了“AI 将取代人类”等消极言论的扩散(Yam et al., 2025)。

因此, 如何有效地建立人类与生成式 AI 之间的区分, 成为顺应智能技术发展、实现人类与 AI 和谐共生过程中亟须解决的关键命题。

然而, 当前研究更多关注人们在追求人类与 AI 区分时表现出的直接对抗行为(Mariadassou et al., 2024)。例如当生成式 AI 的身份被凸显时, 个体往往忽视其真实表现, 而通过贬低或否认其能力来突出 AI 与人类的差异(Hitsuwari et al., 2023; Magni et al., 2024)。但是, 这种“自欺欺人”的方式不仅未能给出生成式 AI 与人类间的真实区分, 甚至还将加剧人类与 AI 间的隔阂(Giger et al., 2024; Reich et al., 2023)。基于此, 本文将探讨生成式 AI 凸显时个体是否存在其他积极方式来建立人类与生成式 AI 的区分, 以及这一过程的调节因素。探讨这一问题不仅有助于深化对人类与 AI“和而不同”共存方式的理解, 更能够为推动人类积极拥抱智能技术的发展带来启发。

收稿日期: 2024-01-30

* 国家社科项目(19F5HB010); 教育部高等学校教指委教改项目(20221039); 教育部产学研合作协同育人项目(231000287155910)资助。周详和白博仁为本文共同第一作者。

通信作者: 周详, E-mail: zhouxian@nankai.edu.cn; 张婧婧, E-mail: 1120210958@mail.nankai.edu.cn

1.1 生成式 AI 凸显威胁人类心智独特性

生成式 AI 凸显模糊了人类心智的独特性界限。生成式 AI 凸显(salience)是指生成式 AI 作为一种社会类别(social category)引起个体注意的现象(许丽颖 等, 2022; Jackson et al., 2020; Taylor & Fiske, 1978)。不同于传统的判别式 AI 往往被人们认知为不具备社会属性、仅能机械执行人类指令的程序(Riemer & Peter, 2024), 如今的生成式 AI 凭借其灵活、强大的能力表现被越来越多地赋予了与人类相似的主体性(agentie), 因此也逐渐以人类外群体的身份在群际层面影响人类的观念与行为(许丽颖 等, 2025; Ng & Lin, 2022; Wang & Peng, 2023)。其中最直接的影响即是生成式 AI 对人类群体独特性的挑战。依据社会认同理论(Social Identity Theory, SIT), 群体独特性是一个群体存在与维持的基础, 内群体成员会积极寻求所在群体的独特性并以此为基础建立群体认同(黄殷, 寇戎, 2013; Tajfel & Turner, 1979)。长久以来, 心智(minds)被视为人类群体独特性的一种具体体现(Turing, 1950; Laland & Seed, 2021), 涵盖感受性维度和能动性维度共 18 项心智能力(Dang & Liu, 2022; Gray et al., 2007)。然而, 生成式 AI 成功模拟了如交流、推理等以往被视为人类独有的心智能力, 动摇了人类心智的独特地位(Santoro & Monin, 2023)。近年来, 诸如“生成式 AI 突破图灵测试、接近人类智能”等话题在媒体广泛传播(Mitchell, 2024), 进一步强化了生成式 AI 作为一种人类外群体却拥有心智的观念。因此, 如今仅仅是生成式 AI 凸显便会提醒人们人类心智独特性正在丧失, 而这种威胁也将促使人们采取行动维护人类心智的独特性。

1.2 生成式 AI 凸显下的社会创造策略

事实上, 研究者已开始关注生成式 AI 凸显下个体如何维护人类心智独特性, 但多聚焦于社会竞争策略。依据社会认同理论, 维护群体独特性存在社会竞争策略(Social Competition Strategy)与社会创造策略(Social Creativity Strategy)两种手段(王沛, 刘峰, 2007; Tajfel & Turner, 1979)。其中, 社会竞争策略是指通过直接竞争的方式来突出群体间的差异, 往往要求内群体成员关注与外群体存在重叠的特征维度(下文简称为重叠维度)。内群体成员需要设法在重叠维度上成为优势一方, 从而实现与劣势外群体的区分(van Bezouw et al., 2021)。已有研究中观察到的个体打压 AI 创造力表现(Millet et al., 2023)、否认 AI 学习能力(Reich et al., 2023)等现象,

均可以被视为个体为了维护人类群体的独特性, 采用贬低、诋毁等直接对抗的方式导向当前 AI 在创造性、学习能力等方面仍劣于人类的结论。然而, 在生成式 AI 技术蓬勃发展的背景下, 社会竞争策略在维护人类群体独特性中的作用愈发有限(Cha et al., 2020), 因此这也驱使研究者将视角转向社会创造策略。

然而, 当前研究在解读社会创造策略的内涵时存在一定局限。Tajfel 和 Turner (1979)提出社会创造策略是指通过重新定义(redefine)群际比较的方式突出群体间差异。这种策略要求内群体成员回避重叠维度, 转而通过比较内群体与外群体尚未重叠的特征维度(下文简称为区分维度)以突出群体间差异(van Bezouw et al., 2021)。例如 Cha 等人(2020)发现, 当 AI 与人类在理性和精细性特征发生重叠时, 人们会转而将文明性、主体性等 AI 尚不具备的特征作为区分维度, 并通过在区分维度上对比人类与 AI 来突出人类的独特性。然而, 已有研究仅将该策略名称中的“创造性”操作化为一种与直接竞争相对的“灵活逃避”(Cha et al., 2020; Douglas et al., 2005; Santoro & Monin, 2023), 忽视了群体成员也可以主动创造群际差异。本文提出, 社会创造策略更重要的内涵在于群体成员会加深对内群体独特性的理解, 主动扩展区分维度内容。具体而言, 当外群体与内群体在某些特征上存在重叠时, 内群体成员并不会一味“割让”这些重叠维度特征并局限于已知的区分维度, 而是会审视并细化这些重叠维度特征, 从中挖掘出新的、仍未被重叠的独特性特征补充至区分维度。以非裔美国人为例, 当其被迫与其他种族使用同一种语言(英语)时, 他们并没有简单接受语言独特性的丧失, 而是主动创造出了拥有独特语法、词汇与语音特征的非裔美国人英语(African American Vernacular English), 重新建立了内群体在语言上的独特性(Winford, 2015)。与之类似, 在生成式 AI 凸显情境下, 人们很可能也会审视那些 AI 已模拟的心智能力, 从中“创造性”地发掘出 AI 尚不具备的子维度能力, 再将其与已知 AI 尚不具备的心智能力合并, 共同作为区分维度以突出人类与生成式 AI 的差异。据此, 本文将探讨生成式 AI 凸显下, 个体的社会创造策略是否会表现出主动扩展区分维度的形式?

在此基础上, 本文认为社会创造策略很可能是一种比社会竞争策略更有效的生成式 AI 凸显应对方式。首先, 社会创造策略往往能够形成更突出的

群际差异,以应对生成式 AI 凸显带来的独特性威胁。当人类与生成式 AI 的比较重心从重叠维度转向区分维度时,二者间差异的表现形式也从“我优你劣”变成了“我有你无”。这种作为特质差异的后者比作为程度差异的前者更显著(Santoro & Monin, 2023)。因此,生成式 AI 凸显时,个体很可能会使用社会创造策略以有效维护群体独特性;更重要的是,社会创造策略可有效避免社会竞争策略引发的负面效应。由于社会竞争策略要求内群体成员在重叠维度与外群体进行直接对抗,这种“你追我赶”的竞争可能会进一步强化人类对生成式 AI 的相似性感知(van Bezouw et al., 2021),从而陷入人类与生成式 AI 间相似与竞争的恶性循环。而社会创造策略则要求内群体成员借助区分维度与外群体建立差异,要求人类专注于发掘与扩展自身独特性。这种“另辟蹊径、各美其美”的观念可能更有利于真正化解生成式 AI 带来的独特性威胁(Becker, 2012),走向人类与 AI 和谐共生的未来。据此,本文将探讨生成式 AI 凸显情境下,个体是否会使用社会创造策略?

1.3 人类中心主义的调节作用

人类中心主义(Anthropocentrism)很可能调节生成式 AI 凸显下个体的社会创造策略使用程度。人类中心主义是指个体认为人类优先于其他物种或人造物的信念(Fortuna et al., 2023)。由于生成式 AI 技术模糊了其与人类间的界限并动摇了人类作为唯一拥有心智群体的优势地位,因此人类中心主义很可能会增强人们采取行动应对生成式 AI 独特性挑战的动机(Cubric, 2020)。Millet 等人(2023)即发现人类中心主义会提升个体的社会竞争策略使用程度,表现为个体更加抗拒 AI 艺术。与之相似,大量研究也观察到人类中心主义会增强个体对 AI 的敌意与排斥(Fortuna et al., 2023; Jiang et al., 2024; Yam et al., 2023),并由此导向了人类中心主义作为一种偏见应当被矫正的结论。然而,值得注意的是,现有研究主要在社会竞争情境下考察人类中心主义的影响,例如要求被试评价 AI 表现或表达 AI 厌恶等。因此,这些研究结果主要反映了人类中心主义会促进社会竞争策略使用。然而本文认为,以往研究中仅为被试提供了社会竞争情境,如果研究情境允许个体采用社会创造策略,结果是否会有所不同?鉴于社会创造策略能够更有效地突出人类与生成式 AI 之间的差异,因此相比低人类中心主义者,高人类中心主义者很可能会在更强动机驱动下

更多地使用社会创造策略,从而利用区分维度有效维护人类心智的独特性。因此,本文将考察生成式 AI 凸显情景下,人类中心主义会如何影响个体的社会创造策略使用情况?

1.4 研究概览

综上,本文将考察生成式 AI 凸显情景下个体的社会创造策略使用情况,以及人类中心主义对这一过程的调节作用。将通过 4 项递进研究验证本文的基本假设:个体会创造性地定义人类心智与生成式 AI 的区分方式,并运用此种社会创造策略应对生成式 AI 凸显,而人类中心主义会增强生成式 AI 凸显情境下个体的社会创造策略使用程度。具体而言:研究 1 通过两个子研究明确个体的社会创造策略表现形式,为后续研究奠定基础;研究 2 首先检验生成式 AI 凸显后个体是否会使用社会创造策略;研究 3 进一步引入人类中心主义这一调节变量,检验其是否会增强生成式 AI 凸显情境下个体的社会创造策略使用程度;研究 4 通过补充社会竞争策略情境作为研究控制条件,再次检验生成式 AI 凸显情境下个体是否仍会使用社会创造策略,以及人类中心主义将如何影响社会创造策略使用程度。

2 研究 1: 生成式 AI 凸显下社会创造策略的表现形式

研究 1 旨在通过两个子研究探讨人们为区分人类心智与生成式 AI 所采用的社会创造策略的具体表现形式,为后续研究中对社会创造策略进行测量提供重要依据。具体而言,本研究关注在生成式 AI 模仿人类心智能力的背景下,个体是否不仅会将经典的 18 项心智能力划分为重叠维度与区分维度,还会进一步审视并细化那些 AI 已模拟的心智能力作为区分维度的补充。其中,研究 1a 首先要求被试判断生成式 AI 对经典的 18 项人类心智能力的模仿情况以获得重叠维度与区分维度内容,同时收集被试创造性提出的、可用于区分人类与生成式 AI 的新心智能力;在此基础上,研究 1b 首先要求被试判断生成式 AI 对经典的 18 项人类心智能力及收集的新心智能力的模仿情况,以确定重叠维度与区分维度的完整内容,随后通过探索性因子分析验证新区分能力是否来自于对重叠维度的扩展。

2.1 研究 1a

2.1.1 调查内容

参考 Cha 等人(2020)的研究,研究 1a 首先要求

被试评估经典的 18 项人类心智能力中, 哪些能力已被生成式 AI 模仿(重叠维度), 哪些仍然是 AI 尚不具备的(区分维度)。这 18 项人类心智能力包括感受饥饿、感到恐惧等感受性维度心智能力, 以及记忆、思考等能动性维度心智能力(Dang & Liu, 2022; Gray et al., 2007)。在调查中, 这 18 项人类心智能力以随机顺序呈现, 被试需依次判断生成式 AI 是否模仿了该能力。其中, 判断为“是”次数显著高于期望值的能力将被归类至重叠维度, 显著低于期望值的能力将被归类至区分维度。随后, 被试需进一步作答: 除了上述能力外, 还有哪些生成式 AI 尚不具备的、仍独属于人类的心智能力。开放性作答结果经整理后将由研究 1b 验证其是否可以纳入区分维度, 以及其是否源自被试对重叠维度的扩展。

2.1.2 调查结果

研究 1a 于 Credamo 平台发放并收集到 50 份调查问卷。首先, 分别计算 18 项心智能力中每一项被判断为“是”的情况, 将频率显著高于期望值(概率期望值为 25 次)的能力归类至重叠维度, 显著低于期望值的能力归类至区分维度(Cha et al., 2020)。其中重叠维度包括自我控制($f = 33$)、道德($f = 40$)、记忆($f = 46$)、识别情绪($f = 46$)、计划($f = 48$)、交流($f = 49$)与思考($f = 49$)共 7 项心智能力, $\chi^2(1)s \geq 5.12$, $ps < 0.025$; 区分维度能力包括饥饿($f = 0$)、恐惧($f = 7$)、疼痛($f = 1$)、愉悦($f = 18$)、愤怒($f = 17$)、欲望($f = 11$)、人格($f = 11$)、知觉意识($f = 10$)、骄傲($f = 16$)、尴尬($f = 10$)、开心($f = 18$)共 11 项心智能力, $\chi^2(1)s \geq 3.92$, $ps < 0.05$ 。由此可见, Gray 等人(2007)提出的 18 项心智能力中, 所有能动性维度的心智能力均被归类至重叠维度中, 而所有感受性维度的心智能力均被归类至区分维度中。

此外, 通过整理被试开放性作答的结果, 删除掉与已有心智能力重复(如沟通能力与交流维度重复)以及无意义的作答(如运动能力不属于心智能力), 合并同类作答(如发明不存在事物的能力和深度创新同属于从无到有的创新)后, 共得到 10 项由被试产出的、可用于区分人类与生成式 AI 的新心智能力标签, 分别为从无到有的创新(后简称创新)、经营社会关系、自我意识、共情、直觉、探究、审美、管理情绪、文化信仰、幽默感。研究 1b 将进一步检验这些由被试给出的新心智能力是否可以有效区分人类与生成式 AI, 以及这些能力与重叠维度能力是否同属于能动性维度, 从而验证研究设想。

2.2 研究 1b

2.2.1 调查内容

研究 1b 除检验研究 1a 中由被试生成的新心智能力是否能够有效区分人类与生成式 AI 外, 还将重点考察这些新能力与能动性维度间的结构关系。为了经由探索性因子分析实现这一目标, 参考 Weisman 等人(2017)的研究, 要求被试围绕给定的某一对象(生成式 AI、计算机、甲壳虫、鱼、婴儿、小狗中的任意一个)是否具备调查问卷中列出的心智能力作 0 到 6 点评分, 其中 0 代表完全不具备、6 代表完全具备。在调查中, 共包括 28 项心智能力, 涵盖 18 项经典的心智能力, 及 10 项研究 1a 中由被试生成的新的心智能力。其中, 被试对 10 项新的心智能力的评分将用于检验这些新能力是否能够有效区分人类与生成式 AI。而全部 28 项心智能力的评分结果将在合并后进行探索性因子分析, 通过检验新收集的区分能力与重叠维度能力是否同属于能动性维度, 从而验证研究设想。

2.2.2 调查结果

研究 1b 通过 Credamo 平台发放问卷共 300 份, 其中生成式 AI、计算机、甲壳虫、鱼、婴儿、小狗 6 种对象每种对应收集问卷 50 份。首先, 使用调查对象为生成式 AI 的问卷数据检验 10 项新的心智能力是否可以有效区分人类心智与生成式 AI。单样本 t 检验结果表明, 10 项新能力得分均显著小于中值 3, $ts \leq -2.04$, $ps \leq 0.025$ 。这表明被试认为生成式 AI 不具备这 10 项能力, 因此这 10 项新能力均可以被纳入至区分维度。随后, 使用调查对象为生成式 AI 的问卷中其余 18 项经典心智能力的数据进行单样本 t 检验, 结果再次证明了研究 1a 中对重叠维度与区分维度的划分: 11 项感受性维度能力得分均显著小于中值 3, $ts \leq -4.34$, $ps \leq 0.001$; 而 7 项能动性维度能力得分均显著大于 3, $ts \geq 3.24$, $ps \leq 0.002$ 。因此, 结合研究 1a 与研究 1b 结果可知, 面对生成式 AI 时人类心智能力中的重叠维度包括 7 项能力, 区分维度包括 21 项能力。后续研究将以这一划分为基础, 将个体面对生成式 AI 凸显时的社会创造策略操作化为个体降低心智能力中 7 项重叠维度的重要性评价、提高 21 项区分维度的重要性评价, 以突出人类心智与生成式 AI 差异。

此外, 为探讨被试生成的新心智能力条目否来自于对重叠维度的扩展, 将以生成式 AI、计算机、甲壳虫、鱼、婴儿、小狗为调查对象的所有问卷合并以进行探索性因子分析, 检验新能力条目与原有

感受性维度及能动性维度的从属关系。因子分析的前提检验结果显示 KMO 值为 0.949, Bartlett 球形检验结果为 $\chi^2(378) = 7173.024, p < .001$, 表明数据适合进行因子分析。采用主成分分析法(PCA)进行因子提取, 依据碎石图确定因子数量为 2 个, 并采用最大方差法(Varimax)进行旋转, 以获得更具理论意义的因子结构, 结果如表 1 所示。其中, 因子 1 解释了 37.32% 的方差, 包括经典的 11 项感受性维度能力以及自我意识和直觉 2 项新的心智能力, 因此将该因子命名为感受性维度。因子 2 揭示了 25.44% 的方差, 包括经典的 7 项能动性维度能力, 以及创新、经营社交关系等 8 项新的心智能力, 因此将该因子命名为能动性维度。因子分析结果显示,

表 1 探索性因子分析结果

心智能力	因子 1 (感受性维度)	因子 2 (能动性维度)
饥饿	0.87	
恐惧	0.88	
疼痛	0.90	
愉悦	0.89	
愤怒	0.86	
欲望	0.83	
个性	0.75	
知觉意识	0.87	
自豪	0.50	0.49
尴尬	0.54	0.52
开心	0.87	
自我控制		0.48
道德		0.82
记忆		0.67
识别情绪		0.59
计划		0.82
交流		0.77
思考		0.84
创新		0.69
经营社会关系		0.61
自我意识	0.83	
共情		0.70
直觉	0.68	
探究	0.46	0.47
审美		0.66
管理情绪	0.48	0.60
文化信仰		0.63
幽默		0.75

注: 表中仅显示大于 0.4 的载荷。其中, 前 11 项(饥饿-开心)为 Gray 等人(2007)提出的感受性维度心智能力, 12~18 项(自我控制-思考)为 Gray 等人(2007)提出的能动性维度心智能力, 19~28 项(创新-幽默)为被试提出的新心智能力。

被试提出的新能力中, 除少数属于感受性维度外, 绝大部分归属于能动性维度。这表明, 在面对生成式 AI 凸显时, 个体并非简单依赖传统的感受性维度建立区分维度, 同时也会通过创造性地细化能动性维度来构建区分维度。该结果支持了研究设想, 即当能动性维度中的能力均被生成式 AI 模仿时, 社会创造策略不仅包括将感受性维度纳入区分维度, 还包括个体“创造性”地从能动性维度中扩展出可纳入区分维度的新心智能力。

3 研究 2: 生成式 AI 凸显引发个体使用社会创造策略

研究 2 将探讨个体是否会使用社会创造策略以应对生成式 AI 凸显。拟验证的研究假设为: 生成式 AI 凸显后, 个体会使用社会创造策略。

3.1 被试

使用 G*Power 3.1 软件(Faul et al., 2007)计算研究所需样本量。统计方式为单因素三水平 F 检验, 取中等效应量 $f = 0.25$, 显著性水平 $\alpha = 0.05$, 为达到 80% 统计检验力需要至少 159 名被试。通过在 Credamo 平台发布实验程序共获得 180 份数据, 剔除没有通过注意检查的被试数据, 最终得到有效数据 177 份, 被试平均年龄为 31.23 ± 8.18 岁, 男性 64 人(占 36.2%)。其中生成式 AI 凸显组被试 61 人, 人类凸显组被试 57 人, 无凸显组被试 59 人。所有被试均为自愿参加实验且知情同意, 实验结束后获得相应实验报酬。实验已获得所在机构伦理审批。

3.2 实验设计与程序

研究 2 为单因素被试间实验设计。自变量为凸显组别(生成式 AI 凸显组/人类凸显组/无凸显空白组), 因变量为社会创造策略使用程度。

所有被试在阅读并签署知情同意书后被随机分配至三组中的任意一组并开始正式实验。首先, 对生成式 AI 凸显进行操纵。参考 Jackson 等人(2020)操纵凸显的形式, 制作两张内容相似的新闻网页截图用于操纵生成式 AI 凸显及人类凸显。依据本文对生成式 AI 凸显的定义, 操纵材料中需要引导被试注意到生成式 AI 这一社会类别。因此, 生成式 AI 凸显组材料中的主体被声明为非特指化的“生成式 AI”, 以引导被试将材料中提及的生成式 AI 视为一种社会类别。与之相对, 人类凸显组材料中主体的社会类别为人类, 操纵材料中具体使用了职场新人这一表述以保证材料内容的合理性与流畅性。除主体的社会类别不同外, 材料中的其他陈述均相

同,具体内容如图 1 所示。作为操纵检验,被试在阅读完材料后需要回答“请问刚才材料中提到的是哪一类群体具备了多种能力”,以确认其是否成功识别了材料中主体的社会类别。无凸显组中不呈现任何材料,直接进入后续实验流程。

随后,参考既往研究(Cha et al., 2020),使用人类心智能力重要性问卷测量被试的社会创造策略使用情况。基于研究 1 结果,人类心智能力重要性问卷共包含 28 个条目,其中重叠维度包括沟通、识别情绪等共 7 项能力(Cronbach's $\alpha = 0.76$),区分维度包括能感受饥饿、从无到有的创新等共 21 项能力(Cronbach's $\alpha = 0.93$)。被试需要思考并回答“下列心智能力对人类来说分别有多重要? ”。问卷使用 7 点评分量表,其中 0 代表完全不重要,6 代表非常重要。在此基础上,基于研究 1 中提出的社会创造策略的操作性定义:“个体降低心智能力中 7 项重叠维度的重要性评价、提高 21 项区分维度的重要性评价,以突出人类心智与生成式 AI 差异”,本研究将计算被试对区分维度重要性平均分与重叠维度重要性平均分的差值,将其作为个体社会创造策略使用程度的测量指标。与 Cha 等人(2020)分别比较重叠维度与区分维度重要性评价的方法相比,使用区分维度与重叠维度的重要性差值一方面能够整合同一被试对两类维度的重要性调整,使后续统计分析更加高效;另一方面也能够减少个体评价习惯带来的误差,提高组间比较的敏感性。因此,本研究通过计算该差值以从整体上概括个体在定义人类特征时是否更倾向于强调人类心

智与生成式 AI 之间的区分维度能力、回避二者间的重叠维度能力。差值越大,代表个体赋予区分维度的相对重要性水平更高,即社会创造策略的使用水平更高。

最后,被试将报告性别、年龄两项人口统计学信息并作答控制变量题项。为避免被试原有接触生成式 AI 的经验与研究中生成式 AI 凸显的效应相混淆,参考 Cha 等人(2020)研究纳入了对生成式 AI 的关注与熟悉程度作为控制变量,共 2 个条目,均采用 5 点计分:(1)请问您是否经常关注生成式 AI 相关的信息和新闻;(2)请问您认为,您对生成式 AI 的了解程度与一般人相比怎么样。完成全部作答任务后主试将发放报酬并解答研究中的疑虑。

3.3 结果

所有纳入统计的生成式 AI 凸显组 61 名被试与人类凸显组 57 名被试均准确回答了材料中主体的身份类别,表明操纵有效。

以组别作为自变量,性别、年龄以及对生成式 AI 的关注和熟悉程度为协变量,社会创造策略使用程度为因变量进行单因素三水平方差分析,结果表明组别主效应显著, $F(1, 170) = 6.31, p = 0.002, \eta_p^2 = 0.07$ 。事后比较发现,生成式 AI 凸显组 ($M = -0.23, SD = 72$) 的社会创造策略使用程度显著高于人类凸显组 ($M = -0.54, SD = 0.43$), $p = 0.017$, 以及无凸显组 ($M = -0.64, SD = 0.78$), $p = 0.001$, 而人类凸显组与无凸显组之间差异不显著, $p = 0.35$ 。

3.4 讨论

研究 2 的结果发现,相较于人类凸显组和无凸

所在位置: 首页 > 深瞳 > 深瞳聚焦 > 正文

职场新人的多样化能力: 推动企业进步的关键力量

社会综合 来源: ABC News 2024-08-19

美国市场调查公司Kadence通过调研市面上多家企业后发现,近年来许多初入职场的新人已经具备了丰富的能力。



例如,他们不仅能够进行顺畅的沟通,还能察觉同事的情绪,提供更具针对性的回应。此外,他们能够记住过往任务的关键信息,在后续工作中调整策略,提升表现。另外,职场新人也可以思考问题,制定计划,帮助企业优化流程,提高效率。一些职场新人还展现出了良好的道德判断能力,能够在工作中遵循职业伦理规范。与此同时,他们也能够很好地管理自身行为,适应工作环境,确保在公司规则范围内发挥最佳表现。

随着经验的生长,刚进入职场的新人们正不断拓展自己的能力,他们在企业中的角色也日益重要。未来从金融到医疗,从客户支持到市场分析,他们或许会成为企业运营的重要组成部分。

所在位置: 首页 > 深瞳 > 深瞳聚焦 > 正文

生成式AI的多样化能力: 推动企业进步的关键力量

社会综合 来源: ABC News 2024-08-19

美国市场调查公司Kadence通过调研市面上多款生成式AI后发现,近年来涌现的生成式AI已经具备了丰富的能力。



例如,生成式AI不仅能够进行顺畅的交流,还能识别用户的情绪,提供更具针对性的回应。此外,它们还能记住过往交互信息,在后续对话中调整回答,提高用户体验。另外,生成式AI也可以思考问题,制定计划,帮助企业优化流程,提高效率。一些生成式AI还能够理解道德观念,从而保证其输出结果符合伦理规范。与此同时,生成式AI还可以控制自身行为,避免非预期的偏差,确保在设定的规则范围内运作。

随着技术的不断进步,生成式AI的能力仍在持续拓展,其在商业世界中的潜力也将愈发不可忽视。未来从金融到医疗,从客户支持到市场分析,它们或许会成为企业运营的重要组成部分。

图1 人类凸显组与生成式 AI 凸显组的操纵材料

显组, 生成式 AI 凸显组的被试应用社会创造策略的程度更高。这一结果表明, 当个体面临生成式 AI 的凸显时, 他们会将人类心智定义的重点向区分维度中的能力倾斜, 来突出人类心智的独特性。因此, 研究 2 对本文的主要假设, 即生成式 AI 凸显后个体会使用社会创造策略, 提供了实证支持。然而, 个体所持的人类中心主义信念是否会影响这一过程仍有待探讨。为此, 研究 3 将进一步检验人类中心主义是否调节生成式 AI 凸显对社会创造策略使用的影响。

4 研究 3: 人类中心主义的调节作用

研究 3 引入人类中心主义作为调节变量, 探讨生成式 AI 凸显下社会创造策略的使用程度是否受到个体人类中心主义水平的影响。拟验证的研究假设为: 人类中心主义会调节生成式 AI 凸显时个体社会创造策略的使用程度。相较于低人类中心主义者, 高人类中心主义者在生成式 AI 凸显的情境下社会创造策略使用程度更高。

4.1 被试

使用 G*Power 3.1 软件计算研究所需样本量。采用线性回归分析检验调节效应, 以 ΔR^2 衡量调节效应大小, 取中等效应量 $f^2 = 0.1$, 显著性水平 $\alpha = 0.05$, 被检验的预测变量为 1 个, 为达到 80% 统计检验力需要至少 81 名被试。通过在 Credamo 平台发布实验程序, 获得有效数据 100 份, 被试平均年龄为 30.41 ± 6.59 岁, 男生 32 人 (占 32%)。其中生成式 AI 凸显组被试 49 人, 人类凸显组被试 51 人。所有被试均为自愿参加实验且知情同意, 实验结束后获得相应实验报酬。实验已获得所在机构伦理审批。

4.2 实验设计与程序

鉴于调节效应分析要求自变量的不同水平具备清晰的操作性区分 (Hayes, 2013), 且研究 2 中并未观察到人类凸显组与无凸显组间的显著差异, 本实验仅保留人类凸显组作为对照组。因此, 研究 3 采用单因素两水平的被试间设计。自变量为凸显组别 (生成式 AI 凸显组/人类凸显组), 因变量为社会创造策略使用程度, 调节变量为人类中心主义。

首先, 被试需要先填写由 Fortuna 等人 (2023) 编制的人类中心主义信念量表作为调节变量测量, 共包含 8 个条目: (1) 人类是自然进化的最后一环, 或者从宗教的角度来说, 是“造物之巅”; (2) 人类是独一无二的, 是宇宙中与众不同的存在; (3) 只有人类才能拥有“自我”和内心世界; (4) “人类独一无二”

只不过是人们编造出来的神话 (反向计分); (5) 只有人类才能客观地认识世界的本来面貌; (6) 人的利益比其他生物的需要更重要; (7) 环保倡导者应该记住, 他们行动的首要目标应该是为了人的利益; (8) 只有人类才能欣赏这个世界的美。该量表使用 7 点评分, 1 代表非常不同意, 7 代表非常同意, 本实验中该量表 Cronbach's $\alpha = 0.91$ 。然后, 要求被试阅读操纵材料。本研究中生成式 AI 凸显与人类凸显的操纵材料及操纵检验题项均同研究 2。紧接着, 与研究 2 相同, 被试将完成人类心智能力重要性问卷作为社会创造策略使用程度测量, 其中区分维度能力 Cronbach's $\alpha = 0.88$, 重叠维度能力 Cronbach's $\alpha = 0.82$ 。最后, 被试需填答人口学信息以及对 AI 的关注和熟悉程度作为控制变量测量。完成全部作答任务后主试将发放报酬并解答研究中的疑虑。

4.3 结果

生成式 AI 凸显组 49 名被试与人类凸显组 51 名被试均准确回答了材料中主体的身份类别, 表明操纵有效。

首先, 以被试组别 (生成式 AI 凸显组 = 2, 人类凸显组 = 1) 为固定因子, 社会创造策略使用程度为因变量, 在控制了性别、年龄以及对 AI 的关注和熟悉程度作为协变量后, 单因素方差分析结果证明组别主效应显著, 生成式 AI 凸显组被试的社会创造策略使用程度 ($M = -0.13$, $SD = 0.68$) 显著高于人类凸显组 ($M = -0.49$, $SD = 0.45$), $F(1, 94) = 10.61$, $p = 0.002$, $\eta_p^2 = 0.10$ 。

在此基础上继续检验人类中心主义的调节作用。使用 Hayes (2013) 提供的 SPSS 插件 PROCESS (Model 1), 设置被试组别作为自变量, 人类中心主义为调节变量, 社会创造策略使用程度作为因变量, 性别、年龄以及对 AI 的关注和熟悉程度作为协变量, 设定 Bootstrap 样本量为 5000, 采用偏差校正方法, 选取 95% 置信区间进行调节效应检验。结果显示, 人类中心主义与组别交互作用显著 ($effect = 0.03$, $SE = 0.01$, $t = 2.70$, $p = 0.008$), 调整后模型 $\Delta R^2 = 0.06$, $F(1, 92) = 7.28$, $p = 0.008$ 。交互作用如图 2 所示, 简单斜率分析结果表明, 对低人类中心主义个体而言, 不同凸显组别对社会创造策略使用程度影响不显著 ($effect_{低} = 0.06$, $SE = 0.17$, $t = 0.35$, $p = 0.73$); 而对中、高人类中心主义个体而言, 相比人类凸显组, 个体在生成式 AI 凸显组下的社会创造策略使用程度显著更高 ($effect_{中} = 0.37$, $SE = 0.11$, $t = 3.31$, $p = 0.001$; $effect_{高} = 0.69$, $SE = 0.16$, $t = 4.31$, $p < 0.001$)。

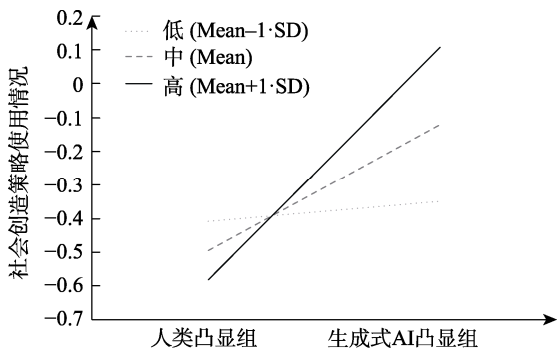


图2 人类中心主义的调节作用

4.4 讨论

研究3不仅再次证明了个体会使用社会创造策略应对生成式AI凸显,还进一步探索了该过程的调节因素。结果发现人类中心主义信念会调节生成式AI凸显时个体社会创造策略的使用程度。具体而言,相较于低人类中心主义者,高人类中心主义者在生成式AI凸显的情境下对区分维度心智能力的重要性评价显著更高,即应用了更高水平的社会创造策略。

然而,考虑到研究2和3均仅设置了社会创造策略情境,与现实生活中人们可以自由选择社会创造策略或社会竞争策略的情况存在差异。如若同时向被试提供社会竞争策略情境,生成式AI凸显后个体是否仍会使用社会创造策略,以及人类中心主义是否仍会提升个体的社会创造策略使用程度尚不清楚。因此,接下来的研究4将聚焦于生成式AI凸显情景并更换研究范式,在同时设置社会竞争策略情境的情况下,考察被试的社会创造策略使用情况以及人类中心主义信念对该策略使用的影响,从而更稳健地检验研究2和3结论在现实生活场景中的有效性与可推广性。

5 研究4:引入社会竞争策略作为控制条件的稳健性检验

为更贴近现实场景,提升研究结论的解释力与外部效度,研究4将同时呈现社会创造策略与社会竞争策略,以测量个体在生成式AI凸显背景下两类策略的实际使用程度。为确保该测量方式能够准确反映社会创造策略与社会竞争策略的概念内涵,所有条目均明确提及生成式AI的存在,故这一测量仅适用于生成式AI凸显情境。且本文核心目的在于检验生成式AI凸显情境下,个体是否倾向使用社会创造策略,以及人类中心主义的影响。因此,研究4将聚焦于生成式AI凸显情境,不设置对照

组。本研究拟验证的研究假设为:(1)生成式AI凸显情景下,个体会使用社会创造策略;(2)人类中心主义会提升个体在生成式AI凸显情境下的社会创造策略使用程度。

5.1 被试

使用G*Power 3.1软件计算研究所需样本量。考虑到实验采用了单样本 t 检验、配对样本 t 检验与线性回归分析,为确保统计检验的充分性,样本量计算基于需要更大样本量的线性回归分析。采用中等效应量 $f^2 = 0.15$,显著性水平 $\alpha = 0.05$,预测变量总数 $= 6$,为达到80%统计检验力需要至少98名被试。通过在Credamo平台发布实验程序发放问卷150份,剔除没有通过注意检查的被试数据,最终得到有效数据145份,被试平均年龄为 31.39 ± 7.14 岁,男生57人(占39.3%)。所有被试均为自愿参加实验且知情同意,实验结束后获得相应实验报酬。实验已获得所在机构伦理审批。

5.2 实验设计与流程

研究4采用横断面相关性设计。自变量为人类中心主义,因变量为社会创造策略使用程度及社会竞争策略使用程度。

首先,被试需要作答与研究3相同的人类中心主义信念量表(Cronbach's $\alpha = 0.90$)。随后,所有被试均将接受与研究2相同的生成式AI凸显操纵并作答操纵检验题项。然后,测量个体对社会创造策略与社会竞争策略的应用情况。考虑到既往研究中社会竞争与社会创造策略的测量方式存在较大差异,难以直接比较,本研究参考Douglas等人(2005)研究中社会创造策略与社会竞争策略的表现形式,围绕人类心智与生成式AI之间的关系,分别生成了8条符合社会创造策略的表述以及与之内容对应的8条符合社会竞争策略的表述,用于测量个体应用两种策略的程度。示例表述为:区分人类心智与AI的关键在于找到那些人类能做到而AI做不到的事(社会创造策略);区分人类心智与AI的关键在于努力让人类做的比AI更好(社会竞争策略)¹。16条表述以随机顺序呈现,要求被试根据认同程度对每条表述作0~6点打分,其中0点代表完全不同意,6点代表完全同意。随后计算被试对社会创造策略8条陈述(Cronbach's $\alpha = 0.74$)与社会竞争策略8条陈述的认同程度(Cronbach's $\alpha = 0.85$)的平均分,分数越高代表个体使用该种策略的程度越高。最后,被

¹ 受限于篇幅限制,全部16条表述请见网络版附录。

试需报告性别和年龄以及对 AI 的关注和熟悉程度作为控制变量测量。完成全部作答任务后主试将发放报酬并解答研究中的疑虑。

5.3 结果

首先, 单样本 t 检验结果表明, 被试在生成式 AI 突显下的社会创造策略使用程度($M = 4.40$, $SD = 0.79$)得分显著高于中值 3(即不确定), $t(144) = 21.39$, $p < 0.001$, Cohen's $d = 1.78$ 。配对样本 t 检验结果进一步表明, 被试的社会创造策略使用程度显著高于其社会竞争策略使用程度($M = 2.84$, $SD = 1.16$), $t(144) = 14.70$, $p < 0.001$, Cohen's $d = 1.22$ 。

其次, 以人类中心主义组别为自变量, 性别、年龄、对 AI 的关注和熟悉程度以及被试的社会竞争策略使用程度为协变量, 社会创造策略使用程度为因变量进行线性回归分析。结果显示, 人类中心主义信念水平能够显著正向预测个体的社会创造策略使用程度(见图 3), $\beta = 0.33$, $t(138) = 4.06$, $p < 0.001$ 。

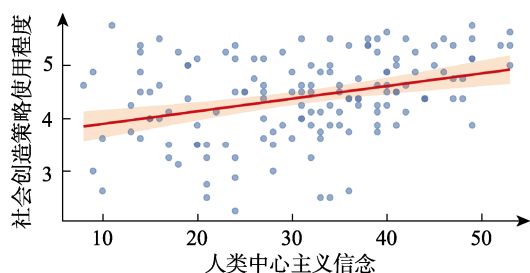


图 3 人类中心主义信念提升个体社会创造策略使用程度

5.4 讨论

研究 4 聚焦于生成式 AI 凸显情景, 通过同时向个体呈现社会创造策略与社会竞争策略, 进一步检验研究 2 和研究 3 的结论是否能够推广至现实场景中。研究 4 结果首先表明, 在社会竞争策略同样可用的情况下, 个体仍倾向于使用社会创造策略。此外, 研究 4 还发现, 在控制了由社会竞争策略使用所产生的影响后, 人类中心主义仍能够显著地提升个体的社会创造策略使用程度。综上, 研究 4 的结果再次支持了研究 2 与研究 3 的结论推断, 进一步提升了其在现实场景中的解释力与外部效度。

6 讨论

本文通过 4 项递进研究探讨了生成式 AI 凸显后个体的社会创造策略使用情况, 以及人类中心主

义对这一过程的调节作用。研究 1 首先明确了生成式 AI 与人类心智间的重叠维度与区分维度内容, 揭示了人们为区分人类心智与生成式 AI 所采用的社会创造策略的具体表现形式; 研究 2 证明了生成式 AI 凸显情境下, 个体会使用社会创造策略, 具体表现为提升人类心智中区分维度的相对重要性; 研究 3 引入人类中心主义作为调节变量, 证明了人类中心主义会增强生成式 AI 凸显后的社会创造策略使用程度; 研究 4 为进一步提升研究结论在现实场景中的解释力与外部效度, 同时向被试呈现社会创造策略与社会竞争策略, 再次证明了生成式 AI 凸显情境下个体会使用社会创造策略, 以及人类中心主义对这一过程起调节作用。接下来对以上主要发现进行讨论。

6.1 社会创造策略: 更积极有效的生成式 AI 凸显应对方式

本文发现在生成式 AI 模仿人类心智能动性维度能力的背景下, 个体不仅会将其尚不具备的感受性维度能力作为人类心智与生成式 AI 的区分维度, 同时也并未完全放弃能动性维度, 而是从中创造出新的、人类独有的心智能力作为区分维度的补充, 从而实现人类心智与生成式 AI 的区分。在此基础上, 本研究证明生成式 AI 凸显后, 个体会使用上述社会创造策略。这一发现不仅支持了已有研究结论(Cha et al., 2020; Santoro & Monin, 2023), 证明人们会使用社会创造策略应对 AI 凸显, 同时也通过扩展社会创造策略的概念内涵, 揭示了其在应对生成式 AI 凸显过程中的双重积极意义。

这种积极意义首先体现在社会创造策略可以扩展与加深人们对人类心智独特性的理解。不同于以往研究仅将社会创造策略视为一种不具有积极意义的灵活逃避(Cha et al., 2020; Douglas et al., 2005), 本文认为社会创造策略对群际独特性的形成起到了重要贡献作用。在应用社会创造策略过程中, 外群体与内群体特征的重叠维度并非仅是仅需要被回避的存在, 同时其也是人们主动创造群际独特性的突破口。研究 1 即发现, 尽管生成式 AI 模拟了传统人类能动性维度中所有心智能力, 但个体并未放弃能动性维度, 而是主动扩展出了从属于能动性维度但生成式 AI 仍不具备的新心智能力。这一过程加深了人们关于“何谓人类心智”的理解, 帮助人们发掘了更多人类独特性内容(Wykowska, 2021)。有趣的是, 尽管研究者较少直接关注社会创造策略, 但会在研究工作中广泛应用这一策略。例

如,当生成式 AI 表现出与人类创造力难以区分的内容生成能力时,研究者不仅没有为了突出人类与 AI 的差异回避谈论创造力,反而为了突出人类创造力的独特性罗列出了人类与生成式 AI 在创造性意图、创造性思维过程等方面的差异,在维护人类独特性的同时加深了对创造力的理解(周详等, 2024; Aru, 2025; Runco, 2023)。综上所述,本文通过拓展社会创造策略的理论内涵,揭示了其在智能技术发展过程中帮助人类重塑自身独特性的重要作用。

这种积极意义还体现在社会创造策略有助于缓解生成式 AI 带来的威胁感,推动人类与 AI 关系的和谐发展。本研究发现生成式 AI 凸显下个体会使用社会创造策略,即将注意力转移至自身的独特性上,更加重视人类心智与生成式 AI 的区分维度。这一发现不同于已有研究中观察到的生成式 AI 凸显后个体会基于社会竞争策略贬低与排斥 AI 的现象(Millet et al., 2023),提示了在该情景中个体也存在另一种更积极的反应。具体而言,当个体采用社会创造策略,通过强调内群体独特特征实现内外群体的区分时,生成式 AI 能够从带来独特性威胁的人类外群体转变为一种对人类心智独特性而言不具威胁的“普通”人类外群体(Santoro & Monin, 2023)。这一角色转变有助于人们更加客观、理性地思考自身与生成式 AI 间的关系。例如 Dang 和 Liu (2022)即发现,那些关注人类心智与 AI 之间差异的群体会更倾向于将 AI 视为合作对象,而关注人类心智与 AI 之间相似性的群体则会更倾向于将 AI 视为竞争对手。因此,本文结果提示未来可以充分发挥社会创造策略在推动人们接纳与拥抱 AI 技术进步中的重要作用。

6.2 人类中心主义:被误解的建设性因素

本文还发现人类中心主义会调节生成式 AI 凸显对个体社会创造策略使用的影响。具体表现为当生成式 AI 凸显时,高人类中心主义组相比低人类中心主义组的社会创造策略使用程度更高,并且在控制社会竞争策略的影响后,人类中心主义仍会增强社会创造策略使用。

这一研究结果首次揭示了人类中心主义对生成式 AI 凸显下个体社会创造策略使用的提升作用。已有研究仅关注了人类中心主义如何调节个体社会竞争策略的使用情况。例如 Millet 等人(2023)发现当生成式 AI 的身份类别被凸显时,高人类中心主义个体的社会竞争策略使用程度更高,具体表现

为其贬低 AI 创造力的程度高于低人类中心主义者。然而本研究发现,当生成式 AI 凸显时,高人类中心主义者也会更多地使用社会创造策略,转而通过强调人类心智中的区分维度以突出人类与生成式 AI 的差异。甚至,在同时提供社会竞争策略情境的情况下,高人类中心主义组仍比低人类中心主义组更多地使用了社会创造策略。尽管这一发现看似与已有研究中将人类中心主义视为一种偏激信念,认为其必然会导致社会竞争策略的观点存在分歧(Millet et al., 2023),然而实则两者之间并不矛盾。作为一种强调人类群体独特性与优越性的信念,人类中心主义的出发点在于提升个体采取行动维护人类独特性的动机(Kim, 2023)。而维护人类独特性的方式除已有研究关注的社会竞争策略外,还存在看似缓和但更有效的社会创造策略(王沛,刘峰, 2007; Tajfel & Turner, 1979)。因此,本文观察到了人类中心主义对个体社会创造策略使用的提升作用。这一发现为已有研究将人类中心主义与贬低、厌恶 AI 等社会竞争策略行为“绑定”的结论提供了补充,更全面地揭示了人类中心主义如何影响个体应对生成式 AI 凸显的方式。

进一步地,本研究还提示了人类中心主义在 AI 时代具有潜在的积极意义。已有研究普遍认为人类中心主义会导向个体排斥与厌恶 AI 的社会竞争策略,因而将其视为一种不利于人类发展进步的、需要被矫正的内群体偏见(Fortuna et al., 2023; Jiang et al., 2024)。而本研究则发现在生成式 AI 凸显情境下,人类中心主义也会提升个体的社会创造策略使用程度。结合前述社会创造策略往往具备着积极影响,此时人类中心主义成为了一种有益于人类发展进步的建设性因素。这一人类中心主义影响方向上的差异可能源自于不同的实验设置下个体会采取不同的应对策略。已有研究仅考虑了社会竞争策略的单一路径,忽视了社会创造策略这一路径的存在,因此也仅观察到了人类中心主义导致个体排斥与厌恶 AI 的消极影响。但若允许被试使用社会创造策略,人类中心主义则可以经由社会创造策略发挥建设性作用。这一发现不仅扩展了已有关于人类中心主义的理论内容,同时也启发了实践中借助社会创造策略以引导人类中心主义创造积极价值。例如可以鼓励人类中心主义者思考人类与生成式 AI 可以分别胜任哪些不同的工作。这种方式不仅能够减少生成式 AI 对人类的独特性威胁,还能够利用人类中心主义者维护人类独特性的高动机,

使其发挥更高水平的创造性以建立更深层次的人类-生成式 AI 合作关系。

6.3 局限与展望

本研究证明了个体在面对生成式 AI 凸显时会采用社会创造策略, 以及人类中心主义在其中起调节作用。然而, 本研究仍存在一定局限性, 为未来研究提供了进一步探索的方向。

首先, 未来研究可进一步拓展社会创造策略的表现形式。依据社会认同理论, 社会创造策略强调对群际比较的重新定义(Tajfel & Turner, 1979)。然而关于重新定义的方式仍存在较大解读空间。例如个体还可能通过重新定义群际比较的对象来实施社会创造策略(van Bezouw et al., 2021)。Dang 和 Liu (2024)即发现人们会拆分人类群体, 通过将使用 AI 的人类划分为外群体的方式维护人类内群体的纯粹性。然而在这一过程中, 内群体成员往往会对“新外群体”赋予负面特征, 以突出自身独特性(Dang & Liu, 2024)。因此, 此种策略可能更多体现了竞争性而非创造性。然而, 未来研究仍可进一步探索是否存在其他社会创造策略的表现形式, 从而更全面、细致地揭示智能技术进步背景下, 人们如何以创造性方式应对人类独特性的挑战。

其次, 未来研究可进一步探讨社会创造策略的积极后效。本研究在丰富社会创造策略具体内涵的基础上, 揭示了其在生成式 AI 凸显情境下的具体表现形式, 强调了该策略在人类重新界定与塑造自身心智独特性中的积极作用, 及其对于构建和谐人智关系过程的潜在价值。然而, 考虑到研究问题的聚焦性, 本研究未直接测量社会创造策略是否能够引导个体真正提升自身的独特心智能力, 以及其是否能够有效化解 AI 焦虑、人智冲突等。因此, 未来研究可以结合现场实验与纵向设计, 进一步验证社会创造策略的积极影响是否能够反映在个体的外显行为层面, 并探究其效果的持续性, 从而更系统地呈现社会创造策略在个体适应新兴智能技术过程中的潜在支持作用。

7 结论

本研究通过 4 项递进研究考察了生成式 AI 凸显情境下个体社会创造策略的使用情况及边界条件。研究结论如下: 第一, 个体会创造性地定义人类心智与生成式 AI 的区别, 以此作为社会创造策略的具体方式; 第二, 生成式 AI 凸显下个体会使用社会创造策略; 第三, 人类中心主义在这一过程

中发挥调节作用, 具体表现为生成式 AI 凸显情境下人类中心主义会增强个体的社会创造策略使用程度。

参 考 文 献

- Aru, J. (2025). Artificial intelligence and the internal processes of creativity. *The Journal of Creative Behavior*, 59(2), e1530. <https://doi.org/10.1002/jocb.1530>
- Banh, L., & Strobel, G. (2023). Generative artificial intelligence. *Electronic Markets*, 33(1), 63.
- Becker, J. C. (2012). The system-stabilizing role of identity management strategies: Social creativity can undermine collective action for social change. *Journal of Personality and Social Psychology*, 103(4), 647–662.
- Cha, Y. J., Baek, S., Ahn, G., Lee, H., Lee, B., Shin, J. E., & Jang, D. (2020). Compensating for the loss of human distinctiveness: The use of social creativity under Human-Machine comparisons. *Computers in Human Behavior*, 103, 80–90.
- Cubric, M. (2020). Drivers, barriers and social considerations for AI adoption in business and management: A tertiary study. *Technology in Society*, 62, 101257.
- 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.
- Dang, J., & Liu, L. (2024). Extended artificial intelligence aversion: People deny humanness to artificial intelligence users. *Journal of Personality and Social Psychology*. Advance online publication. <https://doi.org/10.1037/pspi0000480>
- Douglas, K. M., McGarty, C., Bliuc, A. M., & Lala, G. (2005). Understanding cyberhate: Social competition and social creativity in online white supremacist groups. *Social Science Computer Review*, 23(1), 68–76.
- Faul, F., Erdfelder, E., Lang, A.-G., & Buchner, A. (2007). G*Power 3: A flexible statistical power analysis program for the social, behavioral, and biomedical sciences. *Behavior Research Methods*, 39(2), 175–191.
- Ferrari, F., Paladino, M. P., & Jetten, J. (2016). Blurring human-machine distinctions: Anthropomorphic appearance in social robots as a threat to human distinctiveness. *International Journal of Social Robotics*, 8, 287–302.
- Fortuna, P., Wróblewski, Z., & Gorbaniuk, O. (2023). The structure and correlates of anthropocentrism as a psychological construct. *Current Psychology*, 42(5), 3630–3642.
- Giger, J. C., Piçarra, N., Pochwatko, G., Almeida, N., Almeida, A. S., & Costa, N. (2024). Development of the beliefs in human nature uniqueness scale and its associations with perception of social robots. *Human Behavior and Emerging Technologies*, 2024, 5569587.
- Gray, H. M., Gray, K., & Wegner, D. M. (2007). Dimensions of mind perception. *Science*, 315(5812), 619.
- Hayes, A. F. (2013). *Introduction to mediation, moderation, and conditional process analysis: A regression-based approach*. Guilford Press.
- Hitsuwari, J., Ueda, Y., Yun, W., & Nomura, M. (2023). Does human-AI collaboration lead to more creative art? Aesthetic evaluation of human-made and AI-generated haiku poetry. *Computers in Human Behavior*, 139, 107502.
- Huang, Y., & Kou, Y. (2013). The effect of group distinctiveness on intergroup bias. *Advances in Psychological Science*, 21(4), 732–739.
- [黄殷, 寇彧. (2013). 群体独特性对群际偏差的影响. *心理*

- 科学进展, 21(4), 732–739.]
- Jackson, J. C., Castelo, N., & Gray, K. (2020). Could a rising robot workforce make humans less prejudiced? *American Psychologist*, 75(7), 969–981.
- Jiang, P., Niu, W., Wang, Q., Yuan, R., & Chen, K. (2024). Understanding Users' Acceptance of Artificial Intelligence Applications: A Literature Review. *Behavioral Sciences*, 14(8), 671.
- Kim, M. S. (2023). Meta-narratives on machinic otherness: Beyond anthropocentrism and exoticism. *AI & Society*, 38(4), 1763–1770.
- Laland, K., & Seed, A. (2021). Understanding human cognitive uniqueness. *Annual Review of Psychology*, 72(1), 689–716.
- Magni, F., Park, J., & Chao, M. M. (2024). Humans as creativity gatekeepers: Are we biased against AI creativity? *Journal of Business and Psychology*, 39(3), 643–656.
- Mariadassou, S., Klesse, A. K., & Boegershausen, J. (2024). Averse to what: Consumer aversion to algorithmic labels, but not their outputs? *Current Opinion in Psychology*, 58, 101839.
- Millet, K., Buehler, F., Du, G., & Kokkoris, M. D. (2023). Defending humankind: Anthropocentric bias in the appreciation of AI art. *Computers in Human Behavior*, 143, 107707.
- Mitchell, M. (2024). The Turing Test and our shifting conceptions of intelligence. *Science*, 385(6710), eadq9356.
- Ng, Y. L., & Lin, Z. (2022). Exploring conversation topics in conversational artificial intelligence-based social mediated communities of practice. *Computers in Human Behavior*, 134, 107326.
- Reich, T., Kaju, A., & Maglio, S. J. (2023). How to overcome algorithm aversion: Learning from mistakes. *Journal of Consumer Psychology*, 33(2), 285–302.
- Riemer, K., & Peter, S. (2024). Conceptualizing generative AI as style engines: Application archetypes and implications. *International Journal of Information Management*, 79, 102824.
- Runco, M. A. (2023). AI can only produce artificial creativity. *Journal of Creativity*, 33(3), 100063.
- Santoro, E., & Monin, B. (2023). The AI Effect: People rate distinctively human attributes as more essential to being human after learning about artificial intelligence advances. *Journal of Experimental Social Psychology*, 107, 104464.
- Tajfel, H., & Turner, J. C. (1979). An integrative theory of intergroup conflict. In W. G. Austin & S. Worchel (Eds.), *The social psychology of intergroup relations* (pp. 33–47). Brooks/Cole.
- Taylor, S. E., & Fiske, S. T. (1978). Salience, Attention, and Attribution: Top of the Head Phenomena. *Advances in Experimental Social Psychology*, 11, 249–288.
- Turing, A. M. (1950). Computing machinery and intelligence. *Mind*, 59(236), 433–460.
- van Bezouw, M. J., van Der Toorn, J., & Becker, J. C. (2021). Social creativity: Reviving a social identity approach to social stability. *European Journal of Social Psychology*, 51(2), 409–422.
- Wang, C., & Peng, K. (2023). AI experience predicts identification with humankind. *Behavioral Sciences*, 13(2), 89.
- Wang, P., & Liu, F. (2007). The social identity threat from the perspective of the social identity theory. *Advances in Psychological Science*, 15(5), 822–827.
- [王沛, 刘峰. (2007). 社会认同理论视野下的社会认同威胁. *心理科学进展*, 15(5), 822–827.]
- Weisman, K., Dweck, C. S., & Markman, E. M. (2017). Rethinking people's conceptions of mental life. *Proceedings of the National Academy of Sciences*, 114(43), 11374–11379.
- Winford, D. (2015). The origins of African American vernacular English. In Jennifer B., Lisa J. G., & Sonja L. L. (Eds.), *The Oxford handbook of African American language* (pp. 85–104). Oxford University Press.
- Wykowska, A. (2021). Robots as mirrors of the human mind. *Current Directions in Psychological Science*, 30(1), 34–40.
- Xu, L., Yu, F., Peng, K., & Wang, X. (2022). The influence of robot salience on workplace objectification. *Advances in Psychological Science*, 30(9), 1905–1921.
- [许丽颖, 喻丰, 彭凯平, 王学辉. (2022). 智慧时代的螺丝钉: 机器人凸显对职场物化的影响. *心理科学进展*, 30(9), 1905–1921.]
- Xu, L., Zhao, Y., Yu, F. (2025). Employees adhere less to advice on moral behavior from artificial intelligence supervisors than human. *Acta Psychologica Sinica*, 57(1), 1–23.
- [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(1), 1–23.]
- Yam, K. C., Eng, A., & Gray, K. (2025). Machine replacement: A mind-role fit perspective. *Annual Review of Organizational Psychology and Organizational Behavior*, 12, 239–267.
- Yam, K. C., Tan, T., Jackson, J. C., Shariff, A., & Gray, K. (2023). Cultural differences in people's reactions and applications of robots, algorithms, and artificial intelligence. *Management and Organization Review*, 19(5), 859–875.
- Zhou, X., Zu, C., Cui, Y. X. (2024). *Creativity and Artificial Intelligence*. Xian: Shaanxi Normal University General Publishing House.
- [周详, 祖冲, 崔虞馨. (2024). *创造力与人工智能*. 西安: 陕西师范大学出版总社.]

Unity without uniformity: Humans' social creativity strategy under generative artificial intelligence salience

ZHOU Xiang, BAI Boren, ZHANG Jingjing, LIU Shanrou

*(Department of Social Psychology in School of Sociology and Academy for Advanced Interdisciplinary Studies,
Nankai University, Tianjin 300350, China)*

Abstract

The rapidly advancing artificial intelligence (AI) technology has sparked concerns akin to opening a Pandora's box. In recent years, the remarkable progress of generative AI has further amplified this uncertainty, bringing forth numerous challenges to human society. One such challenge is the symbolic threat to human distinctiveness posed by generative AI. This threat is difficult to avoid and may generate widespread and intricate negative impacts. Therefore, grounded in social identity theory, this study examines individuals' use of social creativity strategies in response to generative AI salience, while also investigating the moderating role of anthropocentrism in this process.

To investigate these effects, this research conducted four progressive studies. Study 1 used two survey rounds to identify mind dimensions perceived as overlapping or distinct between humans and generative AI, revealing the specific forms of social creativity strategies in response to generative AI salience. Study 2 employed a single-factor, three-level between-subjects design to demonstrate that generative AI salience leads individuals to adopt social creativity strategies, particularly by emphasizing the importance of distinguishing dimensions. Study 3 introduced anthropocentrism as a moderating variable using a 2×2 between-subjects design and found that anthropocentrism strengthens the use of social creativity strategies when generative AI is salient. Finally, Study 4 controlled for the effects of social competition strategies by presenting participants with both social creativity and social competition scenarios, further demonstrating that generative AI salience leads individuals to adopt social creativity strategies, and anthropocentrism enhances social creativity strategy use in generative AI salient contexts.

The findings reveal several key insights across four studies. Study 1 identified the specific mind dimensions that individuals perceive as overlapping or distinct between human and generative AI, demonstrating that people actively redefine these distinctions when faced with generative AI salience. Study 2 showed that individuals respond to generative AI salience by employing social creativity strategies, specifically by emphasizing the importance of distinguishing dimensions of the human mind. Study 3 introduced anthropocentrism as a moderating factor and demonstrated that individuals with higher anthropocentric beliefs are more likely to adopt social creativity strategies when generative AI is salient. Finally, Study 4 controlled for the influence of social competition strategies and provided additional evidence that generative AI salience leads individuals to adopt social creativity strategies, and anthropocentrism enhances the use of social creativity strategies.

In summary, this study demonstrates that generative AI salience prompts the application of social creativity strategy, with a moderating role played by individuals' anthropocentrism. The findings reveal the specific manifestations and positive implications of individuals' use of social creativity strategies in contexts where generative AI is made salient. They also elucidate the functional role of anthropocentrism and offer new insights into how social creativity strategies can be leveraged to promote harmonious human-AI coexistence in the era of intelligence.

Keywords generative artificial intelligence, human mind, social creativity strategy, anthropocentrism

附录：研究 4 社会创造策略与社会竞争策略测量材料

社会创造策略

区分人类心智与 AI 的关键在于找到那些人类能做到而 AI 做不到的事。

人类心智的独特性在于比 AI 的维度更丰富。

人类应当通过发扬自身独有的心智能力来回避与 AI 的直接竞争。

为了维持人类的独特性，我们应当发展那些 AI 不具备的独有能力。

人类可以巧妙地区分出人类心智的独特性，无需与 AI 进行直接对抗。

人类与 AI 之间的关系更多是互补而非竞争。

人类应该关注并努力发扬那些尚没有被 AI 模拟的心智能力。

假如 AI 模仿了人类的某种能力，我们就应该寻找其他人类独有的能力。

社会竞争策略

区分人类心智与 AI 的关键在于努力让人类做的比 AI 更好。

人类心智的独特性在于比 AI 的表现更优。

人类应当通过强化自身的心智能力来进行与 AI 的直接竞争。

为了维持人类的独特性，我们应当超越 AI 的表现。

人类需要通过与 AI 进行直接对抗来区分出人类心智的独特性。

人类与 AI 之间的关系更多是竞争而非互补。

人类应该关注并努力提升那些已经被 AI 模拟的心智能力。

假如 AI 模仿了人类的某种能力，我们就要努力在这些能力上表现得比 AI 更好。