

感知机器人威胁对职场物化的影响^{*}

许丽颖^{1†} 王学辉^{2†} 喻丰¹ 彭凯平²

(¹ 武汉大学心理学系, 武汉 430072) (² 清华大学社科学院心理学系, 北京 100084)

摘要 在“工具人”“打工人”“社畜”等流行语风靡职场的当下, 职场物化已待关注。而随着人工智能尤其是机器人在职场中的使用日益增多, 其产生的职场效应也值得探究。本研究旨在探讨在人工智能飞速发展的当今社会, 人们知觉到机器人的威胁是否会产生或加重职场物化现象。8个递进子研究($N = 3422$)探究了感知机器人威胁对职场物化的影响, 并探索其潜在机制和边界条件。结果发现: 第一, 感知到机器人的威胁会增加人们在职场中物化他人的倾向, 并且感知到机器人认同威胁(即对人类自身独特性的威胁)造成的影响更强; 第二, 控制感在感知机器人威胁(主要是认同威胁)影响职场物化中起中介作用, 感知机器人认同威胁越高, 控制感越低, 职场物化越严重; 第三, 补偿控制的另外三种策略, 即加强个人能动性、支持外部能动性以及肯定特定结构, 能够调节感知机器人威胁对职场物化的影响。研究结果揭示了机器人对人际关系的负面影响及其心理机制, 有助于更好地理解、预警与应对机器人负面社会结果。

关键词 职场物化, 感知机器人威胁, 现实威胁, 认同威胁, 补偿控制理论

分类号 B849: C91

1 引言

职场物化(objectification)并不鲜见, “工具人”“打工人”“社畜”等流行语的一度风靡已然说明问题。物化源于康德之描述, 指“将一个人, 即一个有人性的存在, 降低为物的地位(Kant, 1797/1996, p. 209)”。努斯鲍姆(Martha Nussbaum)深化了此概念, 并以工具性、否定自主性、惰性、可替代性、可侵犯性、所有权及否认主观性概括其特征(Nussbaum, 1995, 1999)。其中, 工具性是物化概念的核心(Orehek & Weaverling, 2017), 并被广泛研究(如 Calogero, 2013; Gervais et al., 2014)。究其本质, 物化即将人(他人或自身)视为达成目的的工具, 否认人性, 进而以对待物的方式对待人(许丽颖等, 2022)。物化研究多集中于性物化(如 Calogero, 2013; Saguy et al., 2010; Schmidt & Kistemaker, 2015; Skowronski et al., 2021), 但物化之对象显然并不局限于女性。职场物

化便是其一, 即在工作场所将人视为物的过程和倾向, 主要反映了工作关系中的工具性和对人性的否定(Baldissarri et al., 2017; Belmi & Schroeder, 2021; 许丽颖等, 2022), 是一种负面的人际互动方式和工具理性。职场物化的影响因素一方面与人相关, 如权力(Gruenfeld et al., 2008)等, 另一方面与工作相关, 如工作的重复性等具体特征(Baldissarri et al., 2017)、工作环境的匿名性(Taskin et al., 2019)等。职场物化的测量通常包括显性测量(如问卷; Belmi & Schroeder, 2021)和隐性测量(如 IAT; Baldissarri et al., 2017)两种。职场物化反映了现代工作中的异化现象(Marx, 1844/1964)。自工业革命以来, 生产力急剧发展, 分工劳动成为可能。个体从人变为职场之“工人”, 每个人都变为社会分工之“螺丝钉”, 当人类劳动从本能变为工具, 人由其生产之价值来衡量本身时, 异化产生, 人遂成工具(Marx, 1844/1964)。随着人工智能技术崛起并渗入生活, 机器人甚至一

收稿日期: 2022-07-12

^{*} 国家自然科学基金青年项目(72101132); 国家社科基金青年项目(20CZX059)支持。

[†] 许丽颖和王学辉为共同第一作者

通信作者: 喻丰, E-mail: psychpedia@whu.edu.cn

同进入职场工作,它或以程序形式辅助思维决策、或以实体形态提供现实帮助(张语嫣 等, 2022)。人机关系逐渐变化、人机界限愈渐模糊,这是否会导致人更将他人物化?

若如此,则当引发忧虑。机器人劳动力的发展若抢占人类劳动岗位,则可能让人产生“人可被当作工具”的合理化信念,即物化(如 Pew Research Center, 2017)。机器人主要可以分为两种:工业机器人和服务机器人。工业机器人是指“自动控制的,可重新编程的,多用途的机械手,三个或更多的轴可编程,可以固定的地方或移动用于工业自动化应用”(ISO 8373:2012),服务机器人则是指“为人类或除工业自动化应用外的设备执行有用任务的机器人”(ISO 8373:2012)。本研究中的职场机器人既包括工业机器人也包括服务机器人,即在职场环境中出现的、与人类一同工作的各种实体机器人。职场机器人之广泛使用对人类就业和安全等方面产生威胁,也对人类身份和人类独特性产生威胁(Yogeeswaran et al., 2016)。毕竟机器人作为一种人造物,即使其智能程度未来能比肩甚至超过人类,它也更可能被当作另一实体,而成为人类的外群体。群际威胁理论(Stephan et al., 2016)表明,外群体会让人产生现实威胁(realistic threat)与象征性威胁(symbolic threat)。前者为资源威胁,涉及到对内群体的身体伤害以及对内群体权力、资源和福祉的威胁;而后者是指对内群体认同、价值观和独特性的威胁(Stephan et al., 2016)。现实威胁自不必说,机器人更易管理,成本也更低(Borenstein, 2011)。它既不需要协调复杂人际关系,也不需要支付除购买维护外的工资,甚至其劳动时间还可以更长(McClure, 2018)。机器人对于技术工作与体力劳动的替代大于创造性工作,而较为年长、教育程度较低且学习能力较弱的工人更容易被取代(Borenstein, 2011)。当然,除了对工作岗位的现实抢占,机器人也可能在互动中对人类安全造成担忧而产生现实威胁(Murashov et al., 2016)。更重要的是,机器人由于其能力与外观拟人,造成了认同威胁(即群际威胁理论中的象征性威胁)。机器人的计算与承重等身体能力依然超越人类,造成人类需要认知失调地对待人类的定义性特征问题(喻丰, 2020; 喻丰, 许丽颖, 2018);同时,机器人外观也可拥有和人别无二致的观感,人与拟人化机器人外貌的模糊,不得不经由拟人化过程(许丽颖 等, 2017; 喻丰, 许丽颖, 2020)而造成对人类自身独特性的认同威胁。

机器人进入职场,让人感知到现实威胁和认同威胁,人便可能丧失控制感,而保持和强化控制是人面对威胁时的主要反应(Greenaway et al., 2014)。威胁倾向模型(threat orientation model)已然表明,控制感(perceived control)缺乏是经历威胁的必然结果(Thompson & Schlehofer, 2008)。获得对自身和外部世界的控制是人类基本心理需求,人有控制自身生活的动力(Landau et al., 2015)。一旦控制感受损,则人将力图进行补偿。补偿控制理论(compensatory control theory; Kay et al., 2008; Kay et al., 2009; Landau et al., 2015)认为,人们对于减少控制感的事件和认知会采取补偿策略,意图将控制感拉回基线。然如何补偿?一种策略是寻求并偏好对世界的简单、清晰和一致的解釋。物化便是忽视他人复杂的信念、欲求与心理活动,将他人简单化为具体之物,从而让人简单、清晰而一致地解释他人主观状态和外部客观世界,进而产生控制补偿。研究发现,如果男性感觉自己缺乏影响女性的能力,他们就会更倾向于将女性物化,后续在职场情境中的研究也发现,在实验中那些对自身影响同事的能力产生怀疑的被试更倾向于将同事物化(Landau et al., 2012)。这种控制补偿策略被称为肯定非特定结构(nonspecific structure)的控制补偿策略。基于此可知,机器人劳动力对人产生现实威胁与认同威胁,而这两种感知机器人威胁会导致控制感减弱,为补偿控制感,人们采用肯定非特定结构方式作出职场物化。

当然,肯定非特定结构是相对于控制补偿的其他三种策略而言的,即加强个人能动性(personal agency)、支持外部能动性(external agency)以及肯定特定结构(specific structure) (Landau et al., 2015)。肯定特定结构是指人们倾向于相信做出特定的行为能够可靠地产生预期结果,而此从行为到结果的清晰路径是针对减少感知控制的具体情境的。举例而言,如果人们在感知到机器人威胁时采用肯定特定结构的策略来补偿感知控制,那么人们可能会相信某种与机器人密切相关的“行为-结果”,如他们拔掉机器人的电源(行为)就一定能够控制住机器人(结果)。特定结构即指机器人本身,而物化人类则是肯定非特定结构,即不直指对人类造成威胁的机器人本身。加强个人能动性是指“认为一个人拥有必要的资源来执行产生特定结果或达到特定目的所需的一种或一组行为的信念”(Landau et al., 2015),如自我肯定自身能够控制机器人的资源和能力,认为自己比机器人的能力更强,可以指挥机

机器人做自己安排的工作等。而支持外部能动性则表明个体可以依赖自我以外的系统,如将个体能动性交付于神或者政府,相信外部系统会调动资源来实现个体目的,以此重获掌控。换言之,人们也能在感知到机器人威胁时更加支持政府或者更加信仰神,以此相信政府或神能够控制机器人并保护他们。肯定非特定结构与这三种策略均不相同,它不直接针对个人对自身资源的信念,也不直接针对个人对外部系统的信念,它针对的是社会和物理环境的某些方面,而这些方面不在感知控制减少的具体情境范围内(许丽颖 等, 2022)。但因为后三种控制补偿策略或曰信念依然能够恢复因感知机器人威胁而受损的控制感,因此其应能调节感知机器人威胁导致的职场物化结果。

虽然依社会认同路径而推导的感知机器人威胁会造成职场物化的结果合乎逻辑,且威胁会对人际关系造成负面影响是直觉所知(许丽颖 等, 2022),即威胁会产生负面人际结果,如增加人们对物质的不安全感,则会导致他们更多地感知到移民和外国工人的威胁,从而增加对反移民政策的支持(如 Frey et al., 2018; Im et al., 2019),但于机器人而言依然存在争议。例如也有研究发现,机器人劳动力作为外群体能够凸显人类的共同身份,从而减少对外类群体的偏见(Jackson et al., 2020)。有趣的是,在同一研究中,从 37 个国家所收集的数据表明,过去 42 年自动化速度最快的国家对外群体的外显偏见也增加了,这一影响在一定程度上与不断上升的失业率有关(Jackson et al., 2020)。这也侧面表明,机器人可能会因其对人类产生的失业等威胁而对人际关系造成负面影响。因此,本文对于感知机器人威胁对职场物化影响的探讨,一方面有助于在人工智能飞速发展的大背景下更好地深入理解职场物化这一热点社会问题和物化研究核心议题(许丽颖 等, 2022),另一方面也能够一定程度上回应人工智能以何种方式影响人际关系这一尚存争议的重要话题,并通过对职场中机器人负面影响的关注预警现实问题。

本文分为 3 个递进研究,基本假设是感知机器人威胁会增加人们对同事的职场物化,其中控制补偿(控制感)作为中介机制,而其他补偿策略则会调节其效应。研究 1 用 2 个子研究探讨感知机器人威胁是否会增加职场物化,在现实中和在网络上重复验证该现象的存在,证明这一现象的稳健性,并排除人类共同身份(泛人类主义)的作用。研究 2 包括

3 个子研究,探究感知机器人威胁为何会增加职场物化,即验证控制补偿的中介机制,解释这一现象的发生原因,并通过不同的操作方式验证。研究 3 包括 3 个子研究,基于控制补偿理论所提出的补偿控制策略,将个体能动性、外部能动性、特定结构信念作为调节变量,探讨这一现象存在的边界。

2 研究 1: 感知机器人威胁与职场物化的关系

研究 1 包含两个子研究以证明感知机器人威胁与职场物化的初步关系,其中研究 1a 采用实际机器人公司中的高生态效度样本以量表相关证明其初步关系,而研究 1b 则采用操纵自变量的实验方法考察其因果关系,并排除泛人类主义影响。

2.1 研究 1a: 感知机器人威胁与职场物化的相关

2.1.1 被试

从 3 家知名的机器人公司中实地采集数据。研究共招募 1587 名被试,排除未填完以及未通过注意检查者共 361 人,最终有效被试为 1226 人,其中男性 830 名(67.7%),女性 396 名(32.3%),平均年龄 31.38 ± 7.01 岁。所有被试自愿参加实验并知情同意。

2.1.2 材料

感知机器人威胁。采用 Yogeewaran 等人(2016)编制的感知机器人威胁量表。该量表有两个维度:感知机器人现实威胁和感知机器人认同威胁,每个维度包含 5 个条目。感知机器人现实威胁的条目如下:(1)“我们日常生活中机器人使用的增加正在导致人类失业”;(2)“机器人无法代替人们的工作”;(3)“从长远来看,机器人对人类的安全和福祉构成直接威胁”;(4)“机器人技术的发展会威胁人类的就业和机会”;(5)“机器人在日常生活中的日益普及对人类安全构成了威胁”。其中条目 2 为反向计分。感知机器人认同威胁的条目如下:(1)机器人在日常生活中的广泛应用使我感到困扰,因为它模糊了人类和机器之间的界限;(2)看起来有生命的机器人是令人不安的,因为它们与人类几乎无法区分;(3)技术的最新进步对人类的本质提出了挑战;(4)机器人领域的技术进步正威胁着人类的独特性;(5)机器人开始模糊人类与机器之间的界限。量表为 Likert7 点量表,1 为非常不同意,7 为非常同意。本研究中感知机器人现实威胁维度的 Cronbach's α 系数为 0.80,感知机器人认同威胁维度的 Cronbach's α 系数为 0.84。

职场物化。职场物化的测量基于 Belmi 和

Schroeder (2021)的研究进行修改(将条目中的主体修改为“同事”)。量表包含 7 个条目, 囊括了物化特征的 7 个维度(Nussbaum, 1999), 修改后的具体条目如下:(1)我重不重视某个同事, 主要是看他/她能为我做些什么;(2)我很少关注同事的希望和意愿;(3)我认为同事对我来说是可替代的;(4)我会迫使同事去做我想让他/她做的事;(5)我会把同事当作一个工具来对待;(6)我会关心同事的感受, 因为我真的在乎他/她的人格;(7)哪怕同事不同意我的意见, 我也会继续我自己的计划。其中条目 6 为反向计分。量表为 Likert 7 点量表, 1 为非常不同意, 7 为非常同意。本研究中该量表的 Cronbach's α 系数为 0.71。

为了进一步增强研究结果的稳健性, 研究还加入了以下控制变量:(1)工作场所是否有机器人: 是/否;(2)每天平均与机器人打交道的时长(以小时为单位);(3)对机器人的熟悉程度(Leo & Huh, 2020): 你对机器人有多熟悉?(1 = 一点也不, 5 = 非常)、与普通中国人相比, 你认为你对机器人有多了解?(1 = 一点也不, 5 = 非常);(4)主观社会阶层(Adler et al., 2000)。以及性别、年龄、职位、学历、公司几项人口统计学变量。

2.1.3 共同方法偏差的控制与检验

本研究通过采用匿名数据收集、部分项目反向计分等措施从程序上控制共同方法偏差(周浩, 龙立荣, 2004)。采用 Harman 单因素检验对数据进行共同方法偏差检验, 未旋转的探索性因子分析结果提取出特征根大于 1 的因子共 3 个, 最大因子方差解释率为 31.22% (小于 40%), 因此本研究不存在严重的共同方法偏差。

2.1.4 结果

职场物化与感知机器人现实威胁呈显著正相关($r = 0.15, p < 0.001$), 与感知机器人认同威胁也呈显著正相关($r = 0.18, p < 0.001$)。加入性别、年龄、工作场所是否有机器人、交互时长、机器人熟悉度

和主观社会阶层作为协变量后, 职场物化与感知机器人现实威胁以及感知机器人认同威胁仍然均呈显著正相关($r_{\text{现实}} = 0.16, r_{\text{认同}} = 0.20, ps < 0.001$)。

2.2 研究 1b: 感知机器人威胁对职场物化的影响

2.2.1 被试

通过软件 G*Power 3.1 来确定所需样本量, 根据已有关于机器人劳动力相关研究(Jackson et al., 2020), 取中等效应量 $f = 0.2$, 显著性水平 $\alpha = 0.05$, 单因素两水平被试间设计需要共 266 名被试才能达到 90% 的统计检验力。由于研究职场物化, 研究设置了填写条件为目前在职, 并且考虑到注意检查会排除掉一些被试, 因此共在 Credamo 平台上共计招募了 507 名被试。其中 106 名被试没有通过注意检查, 最终被试量为 401 人。被试平均年龄为 29.41 岁($SD = 5.24$), 女性 201 名(占 50.1%)。

2.2.2 过程

本研究为单因素两水平被试间设计, 被试被随机分配到高感知机器人威胁组(男性 94 人, 女性 96 人)或低感知机器人威胁组(男性 106 人, 女性 105 人)。参考已有研究(Jackson et al., 2020)的操纵材料, 制作了两张不同的新闻网页图片。如图 1, 在高感知机器人威胁组, 被试阅读了一篇名为“机器人: 取代人类劳动力?”的科技新闻, 新闻描述了机器人如何越来越多地抢走人类的工作。在低感知机器人威胁组, 被试阅读了一篇名为“机器人: 只是一时流行?”的科技新闻, 文章描述了机器人永远无法取代大多数人类工作。两个新闻网页的布局、新闻的长度和格式都是相同的。被试看完新闻材料后填写了机器人威胁量表作为操纵检查(测量同研究 1a, 在本研究中, 感知机器人现实威胁维度的 Cronbach's α 系数为 0.85, 感知机器人认同威胁维度的 Cronbach's α 系数为 0.89)。

在实验之前, 我们还提前选择了 3 种不同的办公场景, 并分别请不同的男性和女性在场景中拍摄,



图1 感知机器人威胁的操纵材料

共拍摄了 6 张员工工作照片(见图 2, 已获得授权)。研究者告诉被试, 这是一项“社会知觉研究”, 他们将看到一篇科技新闻和一些照片, 每张照片中有不同的人, 请被试想象他们正在照片所示的环境中与照片中的人一同工作, 即照片中的人是他们的同事。6 张照片呈现的顺序是随机的, 被试看到每张照片后都填写相应的问卷(职场物化的测量同研究 1a, 将对 6 张照片的得分平均后, 本研究中量表的 Cronbach's α 系数为 0.88), 重复这个过程, 直到对 6 张照片后的问题都回答完毕。



图 2 职场物化对象图片材料

接下来, 被试填写泛人类主义测量, 采用 McFarland 等(2012)编制的泛人类主义量表。被试被呈现 6 组图形, 每组图形都有两个圆圈, 但圆圈的重叠程度各不相同。其中一个圆圈代表“自我”, 另一个圆圈则代表“人类”。被试被要求选择最能代表他们如何看待他们自身与全人类之间关系的图形。例如, 选择两个完全独立的圆圈表示被试认为自身与作为一个整体的人类隔绝, 而选择两个完全重叠的圆圈则表示被试与所有其他人类的极度亲密。

之后填写控制变量。与研究 1a 不同, 本研究使用图片材料作为被试想象的工作场景和物化对象, 因此不涉及工作具体属性和工作场所环境的控制变量。参考已有关于职场物化的研究(Belmi & Schroeder, 2021), 主要的控制变量为权力和照片中人物的特点, 具体条目如下: (1)自身权力: “这张

照片让我感到强大”(1 为非常不同意, 7 为非常同意); (2)社会阶层: 你认为照片中的人处于(1 为低社会阶层, 7 为高社会阶层); (3)吸引力: 你认为照片中的这个人的外表有吸引力吗? (1 为一点也没有吸引力, 7 为非常有吸引力); (4)尊重: 一般来说, 你认为人们会尊重照片中的这个人吗? (1 为绝对不尊重, 7 为绝对尊重); (5)物化对象权力: 你认为照片中的这个人有权力吗? (1 为一点也没有权力, 7 为非常有权力); (6)喜欢: 仅仅是看着照片中的这个人, 你有多喜欢他/她? (1 为一点也不喜欢, 7 为非常喜欢)。

接着进行注意检查, 以上量表中穿插了两道注意检查题目, 如“本题请选择非常同意”。此外, 还对操纵材料的内容进行了注意检查, 题目如下: (1)根据研究, 现在机器人占据的工作岗位比例(1 为很大, 2 为很小); (2)根据这项研究作者的说法, 机器人能力和智能的进步所需的时间将比预期(1 为更长, 2 为更短)。

最后填写年龄、性别、教育程度、年收入、职位级别和工作单位类型六项人口统计学信息。

2.2.3 结果

单因素方差分析结果显示, 高感知机器人威胁组的现实威胁($M = 4.69$, $SD = 0.99$)显著高于低感知机器人威胁组($M = 3.10$, $SD = 0.98$), $F(1, 399) = 260.72$, $p < 0.001$, $\eta_p^2 = 0.40$; 高感知机器人威胁组的认同威胁($M = 4.43$, $SD = 1.26$)也显著高于低感知机器人威胁组($M = 3.22$, $SD = 1.19$), $F(1, 399) = 97.31$, $p < 0.001$, $\eta_p^2 = 0.20$, 表明操纵有效。

单因素方差分析结果显示, 高感知机器人威胁组的职场物化($M = 3.54$, $SD = 1.01$)显著高于低感知机器人威胁组($M = 3.32$, $SD = 0.92$), $F(1, 399) = 4.94$, $p = 0.027$, $\eta_p^2 = 0.01$ 。这一结果是稳健的。控制了被试的性别、年龄以及年收入后, 高感知机器人威胁组的职场物化($M = 3.53$, $SD = 1.01$)仍然显著高于低感知机器人威胁组($M = 3.32$, $SD = 0.92$), $F(1, 395) = 5.19$, $p = 0.023$, $\eta_p^2 = 0.01$ 。继续加入权力作为控制变量, 结果表明高感知机器人威胁组的职场物化($M = 3.53$, $SD = 1.01$)仍然显著高于低感知机器人威胁组($M = 3.32$, $SD = 0.92$), $F(1, 394) = 6.02$, $p = 0.015$, $\eta_p^2 = 0.02$ 。再加入对于照片中人物的评价作为控制变量, 结果表明高感知机器人威胁组的职场物化($M = 3.53$, $SD = 1.01$)仍然显著高于低感知机器人威胁组($M = 3.33$, $SD = 0.92$), $F(1, 389) = 4.01$, $p = 0.046$, $\eta_p^2 = 0.01$ 。

对泛人类主义进行检验, 单因素方差分析结果

显示, 高感知机器人威胁组的泛人类主义($M = 4.32$, $SD = 1.14$)与低感知机器人威胁组($M = 4.28$, $SD = 1.11$)无显著差异, $F(1, 399) = 0.10$, $p = 0.748$, $\eta_p^2 < 0.001$ 。并且, 将泛人类主义作为控制变量, 结果表明高感知机器人威胁组的职场物化($M = 3.54$, $SD = 1.01$)仍然显著高于低感知机器人威胁组($M = 3.33$, $SD = 0.92$), $F(1, 389) = 4.88$, $p = 0.028$, $\eta_p^2 = 0.01$ 。

3 研究2: 控制感的中介作用

研究2包含3个子研究, 研究2a以问卷调查的方式探讨控制感在感知机器人威胁两个维度(现实威胁与认同威胁)影响职场物化中的中介作用; 研究2b用实验操纵感知机器人威胁(不区分维度)的方式重复了这一中介作用; 研究2c则分别启动机器人现实威胁与机器人认同威胁, 来考察控制感究竟在哪种威胁条件下起中介作用。

3.1 研究2a: 控制感在感知机器人威胁影响职场物化中的中介作用(问卷研究)

3.1.1 被试

本研究采用G*Power 3.1软件(Faul et al., 2007)确定所需样本量, 对于本实验适用的相关分析, 参考研究1的相关分析结果取较小效应量 $p = 0.16$, 显著性水平 $\alpha = 0.05$, 要达到90%的统计检验力至少需要402名被试。考虑到可能有被试中途退出或未通过注意检查, 因此本研究在两所高校共招募了480名本科生参与此研究, 完成后可获得相应课程学分。被试通过Qualtrics平台填答网络问卷, 并在研究正式开始前阅读了指导语并知情同意。最终共75名被试未完成研究或未通过注意检查, 得到有效数据405份, 其中女性213名(52.6%), 所有有效被试的平均年龄为18.75岁($SD = 1.61$)。

3.1.2 过程

感知机器人威胁(在本研究中, 感知机器人现实威胁维度的Cronbach's α 系数为0.67, 感知机器人认同威胁维度的Cronbach's α 系数为0.81)与职场物化(由于本研究的被试为本科生, 因此该测量前加入以下指导语: “请想象你毕业后进入一家公司工作, 你有许多同事, 请选择你对以下陈述的同意程度”, 本研究中该量表的Cronbach's α 系数为0.71)的测量同研究1a。

控制感的测量采用Lachman和Weaver(1998)编制的控制感问卷的中文修订版(杨沈龙等, 2016)。该问卷包含两个可相加的维度: 掌控感和限制感(反向计分), 分别为4个条目和8个条目, 共12个

条目。掌控感的条目如下: (1)“任何我决心要做的事情, 我几乎都能做到”; (2)“我要是真想做一件事, 一般都能够找到成功的办法”; (3)“我能否得到我想要的东西, 在我的掌控之中”; (4)“我的未来如何, 主要取决于我自己”。限制感的条目如下: (1)“我能做什么和不能做什么多半是由他人决定的”; (2)“对于我生活中很多重要的事情, 我都无法改变”; (3)“在处理生活中的一些问题时, 我常感到无助”; (4)“我生活中所发生的事情常常在我的掌控之外”; (5)“当我想去做某件事情时, 总会受到其他事情的干扰”; (6)“我几乎不能控制发生在我身上的事情”; (7)“我真的没有办法来解决我所有的问题”; (8)“有时我感到在生活中我被别人支使来、支使去”。量表为Likert 7点量表, 1为完全不同意, 7为完全同意。得分越高表明控制感越高。在本研究中, 全部条目的Cronbach's α 系数为0.83。测量条目中穿插有注意检查题目, 最后填写人口统计学信息。

3.1.3 共同方法偏差的控制与检验

本研究通过采用匿名数据收集、部分项目反向计分等措施从程序上控制共同方法偏差(周浩, 龙立荣, 2004)。采用Harman单因素检验对数据进行共同方法偏差检验, 未旋转的探索性因子分析结果提取出特征根大于1的因子共8个, 最大因子方差解释率为16.75% (小于40%), 因此本研究不存在严重的共同方法偏差。

3.1.4 结果

相关分析结果表明, 职场物化与感知机器人威胁及其两个维度均呈显著正相关($r_{现实} = 0.12$, $p = 0.017$; $r_{认同} = 0.18$, $p < 0.001$), 职场物化与控制感及其两个维度均呈显著负相关($r_{控制} = -0.12$, $p = 0.020$; $r_{掌控} = -0.12$, $r_{限制} = -0.23$, $p < 0.001$)。感知机器人认同威胁与控制感呈显著负相关($r = -0.11$, $p = 0.025$), 而感知机器人现实威胁与控制感负相关不显著($r = -0.07$, $p = 0.152$)。

为了验证控制感是否在感知机器人认同威胁对职场物化的影响中起中介作用, 使用SPSS的PROCESS程序进行偏差校正的Bootstrap检验(Hayes, 2013), 选择模型4, 反复抽样5000次。在95%的置信区间下, 将感知机器人认同威胁作为自变量, 将职场物化作为因变量, 将控制感作为中介变量, 做中介效应分析。中介效应检验结果显示, 95%的Bootstrap置信区间不包含0, 因此控制感中介了感知机器人认同威胁对职场物化的影响(间接效应 = 0.02, $SE = 0.01$, 95% CI [0.002, 0.038]), 在

控制了控制感后,感知机器人认同威胁对职场物化的直接影响仍然显著(直接效应 = 0.11, $SE = 0.03$, 95% CI [0.042, 0.170]),说明控制感起部分中介作用。

3.2 研究 2b: 控制感在感知机器人威胁影响职场物化中的中介作用(实验研究)

3.2.1 被试

本研究通过软件 G*Power 3.1 来确定所需样本量,根据已有关于机器人劳动力相关研究(Jackson et al., 2020),取中等效应量 $f = 0.2$,显著性水平 $\alpha = 0.05$,单因素两水平被试间设计需要共 266 名被试才能达到 90% 的统计检验力。由于研究职场物化,因此设置了填写条件为目前在职,并且考虑到注意检查会排除掉一些被试,在 Credamo 平台上共计招募了 359 名被试。其中 62 名被试没有通过注意检查,最终被试量为 297 人。被试平均年龄为 31.52 岁($SD = 5.68$),女性 184 名(占 62.0%)。

3.2.2 过程

本研究为单因素两水平被试间设计,被试被随机分配到高感知机器人威胁组或低感知机器人威胁组,在最终有效被试中,高感知机器人威胁组 148 人,低感知机器人威胁组 149 人。被试所阅读的操纵感知机器人威胁程度的实验材料同研究 1b,被试看完材料后填写了机器人威胁量表(感知机器人威胁的测量同研究 1a,在本实验中,现实威胁维度的 Cronbach's α 系数为 0.92,认同威胁维度的 Cronbach's α 系数为 0.92)作为操纵检查。然后进行控制感测量(测量工具同研究 2a,在本实验中,控制感全部条目的 Cronbach's α 系数为 0.91)和职场物化测量(职场物化的测量同研究 1a,本实验中该量表的 Cronbach's α 系数为 0.80)。测量条目中穿插有注意检查题目,最后填写人口统计学信息。

3.2.3 结果

单因素方差分析结果显示,高感知机器人威胁组的感知机器人现实威胁($M = 4.83$, $SD = 1.11$)显著高于低感知机器人威胁组($M = 2.83$, $SD = 1.14$), $F(1, 295) = 210.81$, $p < 0.001$, $\eta_p^2 = 0.42$;高感知机器人威胁组的感知机器人认同威胁($M = 4.49$, $SD = 1.37$)也显著高于低感知机器人威胁组($M = 2.91$, $SD = 1.25$), $F(1, 295) = 108.42$, $p < 0.001$, $\eta_p^2 = 0.27$,表明操纵有效。

单因素方差分析结果显示,高感知机器人威胁组的职场物化($M = 2.85$, $SD = 0.90$)显著高于低感知机器人威胁组($M = 2.64$, $SD = 0.65$), $F(1, 295) = 5.49$, $p = 0.020$, $\eta_p^2 = 0.02$ 。相关分析结果表明,感

知机器人威胁程度(低威胁组 = 0, 高威胁组 = 1)与控制感呈显著负相关($r = -0.14$, $p = 0.018$),与职场物化呈显著正相关($r = 0.14$, $p = 0.020$),控制感与职场物化呈显著负相关($r_{控制} = -0.50$, $p < 0.001$)。

为了验证控制感是否在感知机器人威胁对职场物化的影响中起中介作用,使用 SPSS 的 PROCESS 程序进行偏差校正的 Bootstrap 检验(Hayes, 2013),选择模型 4,反复抽样 5000 次。在 95% 的置信区间下,将感知机器人威胁作为自变量(低感知机器人威胁组 = 0, 高感知机器人威胁组 = 1),将职场物化作为因变量,将控制感作为中介变量,做中介效应分析。中介效应检验结果显示,95% 的置信区间不包含 0,因此控制感中介了感知机器人威胁对职场物化的影响(间接效应 = 0.11, $SE = 0.05$, 95% CI [0.020, 0.228]),在控制了控制感后,感知机器人威胁对职场物化的直接影响不再显著(直接效应 = 0.11, $SE = 0.08$, 95% CI [-0.051, 0.266]),说明控制感起完全中介作用。

3.3 研究 2c: 控制感在不同威胁影响职场物化中的中介作用(实验研究)

3.3.1 被试

本研究通过软件 G*Power 3.1 来确定所需样本量,根据已有关于机器人劳动力相关研究(Jackson et al., 2020),取中等效应量 $f = 0.2$,显著性水平 $\alpha = 0.05$,单因素三水平被试间设计需要共 321 名被试才能达到 90% 的统计检验力。由于研究职场物化,因此设置了填写条件为目前在职,并且考虑到注意检查会排除掉一些被试,在 Credamo 平台上共计招募了 359 名被试。其中 10 名被试没有通过注意检查,最终被试量为 349 人。被试平均年龄为 31.49 岁($SD = 7.14$),女性 196 名(占 56.2%)。

3.3.2 过程

本研究为单因素三水平被试间设计,被试被随机分配到现实威胁组、认同威胁组和对照组,在最终有效被试中,现实威胁组 115 人,认同威胁组 117 人,对照组 117 人。

为了对感知机器人威胁程度进行操纵,被试首先需要根据一段指导语进行写作,其中现实威胁组的指导语如下:

“随着人工智能的飞速发展,如今所有行业超过 50% 的工作都已经被机器人所占据,未来机器人将必然作为新的工作者与人类进行竞争,对人类产生了现实生活中的威胁。在不久的将来,机器人还会胜任更多高层次工作,进入人类工作的各个领域,

最终取代人们现在从事的许多工作, 对人类造成更大的现实威胁。请想象并描述你目前的工作岗位现在或以后如何被机器人所威胁和取代, 以及被威胁和取代之后你将会面临的情况(不少于 100 字)。”

认同威胁组的指导语如下:

“随着人工智能的飞速发展, 机器人与人类的相似度越来越高, 人与机器会变得难以区分。在外形方面, 许多机器人与人类的外形相似度极高, 看起来跟真人一模一样, 甚至难以区分; 在能力方面, 许多机器人能力与人类相当甚至超过人类, 对人类的独特性产生了威胁。人工智能技术也会将人和机器进行结合, 在未来可能很难定义什么是人、什么是机器人。请想象并描述一个与人类高度相似甚至无法区分的机器人会对你作为人类一员的独特性所产生的威胁, 以及被威胁之后你将会面临的情况(不少于 100 字)。”

对照组的指导语如下:

“随着人工智能的发展, 机器将会协助人类进行许多社会活动。请想象并描述一个未来机器人协助人类的场景(不少于 100 字)。”

在完成写作任务后, 被试依次填写感知机器人威胁量表(同研究 1a, 在本实验中, 现实威胁维度的 Cronbach's α 系数为 0.82, 认同威胁维度的 Cronbach's α 系数为 0.88)、控制感问卷(同研究 2a, 在本实验中, 控制感全部条目的 Cronbach's α 系数为 0.90)和职场物化量表(同研究 1a, 在本实验中该量表的 Cronbach's α 系数为 0.77), 量表中穿插有注意检查题项如“本题请选择非常不同意”, 最后填写人口统计学信息。

3.3.3 结果

以组别(对照组 = 1, 现实威胁组 = 2, 认同威胁组 = 3)作为自变量, 以感知机器人现实威胁作为因变量进行单因素方差分析, 结果发现组别的主效应显著, $F(2, 346) = 44.35, p < 0.001, \eta_p^2 = 0.20$ 。事后多重比较(Bonferroni)表明, 对照组的感知机器人现实威胁评分($M = 3.70, SD = 0.99, 95\% CI [3.52, 3.88]$)显著低于现实威胁组($M = 4.55, SD = 0.96, 95\% CI [4.38, 4.73]$)和认同威胁组($M = 4.85, SD = 0.95, 95\% CI [4.68, 5.03]$), $ps < 0.001$, 而现实威胁组和认同威胁组的差异不显著, $p = 0.057$ 。以组别作为自变量, 以感知机器人认同威胁作为因变量进行单因素方差分析, 结果发现组别的主效应显著, $F(2, 346) = 52.35, p < 0.001, \eta_p^2 = 0.23$ 。事后多重比较(Bonferroni)表明, 对照组的感知机器人认

同威胁评分($M = 3.54, SD = 1.26, 95\% CI [3.31, 3.77]$)显著低于现实威胁组($M = 4.47, SD = 1.10, 95\% CI [4.27, 4.68]$)和认同威胁组($M = 5.07, SD = 1.10, 95\% CI [4.87, 5.27]$), $ps < 0.001$; 并且, 认同威胁组的感知机器人认同威胁评分显著高于现实威胁组, $p < 0.001$ 。这表明, 操纵效果基本符合预期, 其中认同威胁的操纵对于感知机器人威胁的增加效果较强。

以感知机器人威胁(即组别; 对照组 = 1, 现实威胁组 = 2, 认同威胁组 = 3)作为自变量, 以职场物化作为因变量进行单因素方差分析, 结果发现感知机器人威胁的主效应显著, $F(2, 346) = 3.68, p = 0.026, \eta_p^2 = 0.02$ 。事后多重比较(Bonferroni)表明, 认同威胁组的职场物化评分($M = 3.11, SD = 0.82, 95\% CI [2.96, 3.26]$)显著高于对照组($M = 2.85, SD = 0.72, 95\% CI [2.72, 2.98]$), $p = 0.028$, 而现实威胁组与对照组以及认同威胁组均无显著差异, $ps > 0.05$ 。

为了验证控制感是否在感知机器人威胁对职场物化的影响中起中介作用, 使用 SPSS 的 PROCESS 程序进行偏差校正的 Bootstrap 检验(Hayes, 2013), 选择模型 4, 反复抽样 5000 次。在 95% 的置信区间下, 以感知机器人威胁作为自变量(对照组 = 1, 现实威胁组 = 2, 认同威胁组 = 3)并将其在程序中设置为分类变量, 以职场物化作为因变量, 以控制感作为中介变量, 做中介效应分析。结果表明, 在以对照组为参照时, 现实威胁组通过控制感对职场物化的中介效应值为 0.064, 95% 的 Bootstrap 置信区间为 $[-0.025, 0.166]$, 包含 0, 表明中介效应不显著; 认同威胁组通过控制感对职场物化的中介值为 0.116, 95% 的 Bootstrap 置信区间为 $[0.027, 0.215]$, 不包含 0, 表明中介效应显著; 且加入中介变量控制感后, 认同威胁对职场物化的直接效应为 0.132, 95% 的 Bootstrap 置信区间为 $[-0.030, 0.330]$, 包含 0, 表明其直接效应不再显著, 控制感起完全中介作用。以上结果表明, 感知机器人认同威胁会增加职场物化, 这是由于感知机器人认同威胁会降低人们的控制感, 进而使人更倾向于物化职场中的他人。

4 研究 3: 补偿控制策略作为感知机器人威胁与职场物化关系的调节

研究 3 包括 3 个平行子研究, 分别考察个体能动性、外部能动性、特定结构信念是否能够作为调节变量影响感知机器人威胁对职场物化的作用。3 个子研究均采用实验操纵方式进行, 其中研究 3a 将感知机器人威胁的操纵分为了现实威胁与认同

威胁, 结合研究 2 中现实威胁的不显著结果和研究 3a 中现实威胁的阴性结果, 在研究 3b 与研究 3c 在考察调节时均只采用认同威胁与对照组对比。

4.1 研究 3a: 个人能动性的调节作用

4.1.1 被试

采用软件 G*Power 3.1 来确定本研究所需的样本量, 根据已有关于机器人劳动力相关研究(Jackson et al., 2020), 取中等效应量 $f=0.2$, 显著性水平 $\alpha=0.05$, 单因素三水平被试间设计需要共 321 名被试才能达到 90% 的统计检验力。由于研究职场物化, 因此设置了填写条件为目前在职, 并且考虑到可能有被试无法通过注意检查, 并在 Credamo 平台上共计招募了 340 名被试。其中 10 名被试没有通过注意检查, 最终被试量为 330 人。被试平均年龄为 31.56 岁($SD=7.25$), 女性 179 名(占 54.2%)。

4.1.2 过程

本研究为单因素三水平被试间设计, 被试被随机分配到现实威胁组、认同威胁组和对照组, 在最终有效被试中, 现实威胁组 110 人, 认同威胁组 110 人, 对照组 110 人。对现实威胁组、认同威胁组和对照组的操纵过程同研究 2c。在完成写作任务后, 被试依次填写感知机器人威胁量表(同研究 1a, 在本实验中, 现实威胁维度的 Cronbach's α 系数为 0.81, 认同威胁维度的 Cronbach's α 系数为 0.90)、控制感问卷(同研究 2a, 在本实验中, 控制感全部条目的 Cronbach's α 系数为 0.91)、个人能动性量表和职场物化量表(同研究 1a, 在本实验中该量表的 Cronbach's α 系数为 0.81)。

个人能动性的测量一部分改编自以往的补偿控制文献(Beck et al., 2020; Shepherd & Kay, 2018), 并根据本研究的目的和情境进行了相应修改。这部分测量共包括 3 个条目: (1)我认为我自己可以决定对机器人的控制; (2)我认为我对控制机器人有责任; (3)我认为控制机器人在我的掌控之中。另外, 还根据补偿控制理论中个人能动性的定义, 即“认为一个人拥有必要的资源来执行产生特定结果或达到特定目的所需的一种或一组行为的信念”(Landau et al., 2015), 补充编制了 3 个条目: (1)我认为我可以控制机器人; (2)我认为我有能力可以控制机器人; (3)我认为我有资源可以控制机器人。以上测量个人能动性的条目均为 Likert7 点量表, 1 为非常不同意, 7 为非常同意。本研究中该量表的 Cronbach's α 系数为 0.91。量表中穿插有注意检查题项如“本题请选择非常不同意”, 最后填写人口统计学信息。

4.1.3 结果

以组别(对照组 = 1, 现实威胁组 = 2, 认同威胁组 = 3)作为自变量, 以感知机器人现实威胁作为因变量进行单因素方差分析, 结果发现组别的主效应显著, $F(2, 327) = 22.44, p < 0.001, \eta_p^2 = 0.12$ 。事后多重比较(Bonferroni)表明, 对照组的感知机器人现实威胁评分($M = 3.93, SD = 1.04, 95\% CI [3.73, 4.12]$)显著低于现实威胁组($M = 4.71, SD = 1.04, 95\% CI [4.51, 4.90]$)和认同威胁组($M = 4.74, SD = 0.98, 95\% CI [4.55, 4.92]$), $ps < 0.001$, 而现实威胁组和认同威胁组的差异不显著, $p = 1$ 。以组别作为自变量, 以感知机器人认同威胁作为因变量进行单因素方差分析, 结果发现组别的主效应显著, $F(2, 327) = 22.06, p < 0.001, \eta_p^2 = 0.12$ 。事后多重比较(Bonferroni)表明, 对照组的感知机器人认同威胁评分($M = 3.83, SD = 1.43, 95\% CI [3.56, 4.10]$)显著低于现实威胁组($M = 4.55, SD = 1.28, 95\% CI [4.31, 4.79]$)和认同威胁组($M = 5.00, SD = 1.24, 95\% CI [4.76, 5.23]$), $ps < 0.001$; 并且, 认同威胁组的感知机器人认同威胁评分显著高于现实威胁组, $p = 0.037$ 。这表明操纵效果基本符合预期, 其中认同威胁的操纵对于感知机器人威胁的增加效果较强。

以感知机器人威胁(即组别; 对照组 = 1, 现实威胁组 = 2, 认同威胁组 = 3)作为自变量, 以职场物化作为因变量进行单因素方差分析, 结果发现感知机器人威胁的主效应显著, $F(2, 327) = 4.76, p = 0.009, \eta_p^2 = 0.03$ 。事后多重比较(Bonferroni)表明, 认同威胁组的职场物化评分($M = 3.23, SD = 0.93, 95\% CI [3.06, 3.41]$)显著高于对照组($M = 2.88, SD = 0.71, 95\% CI [2.74, 3.02]$), $p = 0.007$, 而现实威胁组与对照组以及认同威胁组均无显著差异, $ps > 0.050$ 。

为了验证控制感是否在感知机器人威胁对职场物化的影响中起中介作用, 使用 SPSS 的 PROCESS 程序进行偏差校正的 Bootstrap 检验(Hayes, 2013), 选择模型 4, 反复抽样 5000 次。在 95% 的置信区间下, 以感知机器人威胁作为自变量(对照组 = 1, 现实威胁组 = 2, 认同威胁组 = 3)并将其在程序中设置为分类变量, 以职场物化作为因变量, 以控制感作为中介变量, 做中介效应分析。在以对照组为参照时, 现实威胁组通过控制感对职场物化的中介效应值为 0.107, 95% 的 Bootstrap 置信区间为 $[-0.001, 0.244]$, 包含 0, 表明中介效应不显著; 认同威胁组通过控制感对职场物化的中介值为 0.133, 95% 的 Bootstrap 置信区间为 $[0.024, 0.268]$, 不包含 0, 表

明中介效应显著;且加入中介变量控制感后,认同威胁对职场物化的直接效应为 0.217, 95%的 Bootstrap 置信区间为[0.021, 0.412], 不包含 0, 表明其直接效应仍然显著, 控制感起部分中介作用。以上结果表明, 感知机器人认同威胁会增加职场物化, 这是由于感知机器人认同威胁会降低人们的控制感, 进而使人更倾向于物化职场中的他人。

为了验证个人能动性是否在感知机器人威胁对职场物化的影响中起调节作用, 使用 SPSS 的 PROCESS 程序进行偏差校正的 Bootstrap 检验 (Hayes, 2013), 选择模型 1, 反复抽样 5000 次。在 95%的置信区间下, 以感知机器人威胁作为自变量 (对照组 = 1, 现实威胁组 = 2, 认同威胁组 = 3) 并将其在程序中设置为分类变量, 以职场物化作为因变量, 以个人能动性作为调节变量, 做调节效应分析。调节分析结果表明, 在以对照组为参照时, 现实威胁与个人能动性的交互作用不显著 ($b = -0.14$, $SE = 0.11$, $t = -1.30$, $p = 0.193$), 认同威胁与个人能动性的交互作用显著 ($b = -0.32$, $SE = 0.11$, $t = -2.79$, $p = 0.005$), 现实威胁组与对照组的职场物化无显著差异 ($b = 0.84$, $SE = 0.58$, $t = 1.45$, $p = 0.148$), 认同威胁组的职场物化显著高于对照组 ($b = 1.81$, $SE = 0.59$, $t = 3.09$, $p = 0.002$), 个人能动性的高低对职场物化无显著影响 ($b = -0.07$, $SE = 0.09$, $t = -0.75$, $p = 0.452$), 模型的 $\Delta R^2 = 0.02$, $F(2, 324) = 4.02$, $p = 0.019$ 。在低个人能动性条件下, 现实威胁对职场物化的影响不显著 ($b = 0.28$, $SE = 0.17$, $t = 1.58$, $p = 0.116$), 而认同威胁对职场物化有显著影响 ($b = 0.57$, $SE = 0.17$, $t = 3.30$, $p = 0.001$); 在高个人能动性条件下, 现实威胁对职场物化的影响不显著 ($b = -0.03$, $SE = 0.15$, $t = -0.19$, $p = 0.848$), 认同威胁对职场物化的影响也不显著 ($b = -0.10$, $SE = 0.16$, $t = -0.62$, $p = 0.536$)。以上结果表明, 认同威胁与个人能动性有显著的交互作用, 当个人能动性低时, 认同威胁会显著增加职场物化; 而当个人能动性高时, 认同威胁对职场物化影响不显著。由于研究 2 和研究 3a 中现实威胁的效应并不显著, 因此我们将后续研究的推进聚焦于认同威胁, 即对认同威胁组与对照组进行比较。

4.2 研究 3b: 外部能动性的调节作用

4.2.1 被试

采用 G*Power 3.1 软件 (Faul et al., 2007) 确定所需样本量, 对于本实验适用的单因素方差分析, 根据已有关于机器人劳动力相关研究 (Jackson et al.,

2020), 取中等效应量 $f = 0.2$, 显著性水平 $\alpha = 0.05$, 单因素两水平被试间设计需要共 200 名被试才能达到 80% ($1-\beta$) 的统计检验力。考虑到可能有被试中途退出或未通过注意检查, 因此在某高校共招募了 256 名本科生参与此研究, 完成后可获得相应课程学分。被试通过 Qualtrics 平台填答网络问卷, 并在研究正式开始前阅读了指导语并知情同意。最终共 42 名被试未完成研究或未通过注意检查, 得到有效数据 214 份, 其中女性 121 名 (56.5%), 所有有效被试的平均年龄为 20.17 岁 ($SD = 1.37$)。

4.2.2 过程

本研究为单因素两水平被试间设计, 被试被随机分配到认同威胁组和对照组, 在最终有效被试中, 认同威胁组 107 人, 对照组 107 人。在完成写作任务后, 被试依次填写感知机器人威胁量表 (同研究 1a, 在本实验中, 现实威胁维度的 Cronbach's α 系数为 0.61, 认同威胁维度的 Cronbach's α 系数为 0.80)、控制感问卷 (同研究 2a, 在本实验中, 控制感全部条目的 Cronbach's α 系数为 0.76)、外部能动性量表和职场物化量表 (同研究 1a, 在本实验中该量表的 Cronbach's α 系数为 0.66)。

外部能动性的测量一部分改编自以往的补偿控制文献 (Beck et al., 2020; Shepherd & Kay, 2018), 并根据本研究的目的和情境进行了相应修改。这部分测量共包括 5 个条目: (1) 我希望将对机器人的控制权交给政府; (2) 我希望政府代表我控制机器人; (3) 我对机器人的控制似乎由政府决定; (4) 我对机器人的控制由政府负责; (5) 我对机器人的控制在政府的掌控之中。另外, 还根据补偿控制理论中外部能动性的说明, 如采用这一策略的人会放弃对自己生活的自主控制, 将个人能动性交给外部系统如神或政府, 通过对外部系统的依赖来恢复控制感, 相信外部系统会调动资源来实现符合他/她利益的结果 (Landau et al., 2015), 补充编制了 3 个条目: (1) 我认为政府可以控制机器人; (2) 我认为政府有能力可以控制机器人; (3) 我认为政府有资源可以控制机器人。以上测量外部能动性的条目均为 Likert 7 点量表, 1 为非常不同意, 7 为非常同意。本研究中该量表的 Cronbach's α 系数为 0.82。量表中穿插有注意检查题项如“本题请选择非常不同意”, 最后填写人口统计学信息。

4.2.3 结果

单因素方差分析结果显示, 认同威胁组的现实威胁评分 ($M = 4.19$, $SD = 0.78$, 95% CI [4.04, 4.33])

显著高于对照组($M = 3.88$, $SD = 0.71$, 95% CI [3.74, 4.01]), $F(1, 212) = 9.22$, $p = 0.003$, $\eta_p^2 = 0.04$; 认同威胁组的认同威胁评分($M = 4.20$, $SD = 1.18$, 95% CI [3.97, 4.42])也显著高于对照组($M = 3.69$, $SD = 0.95$, 95% CI [3.51, 3.87]), $F(1, 212) = 11.91$, $p = 0.001$, $\eta_p^2 = 0.05$ 。操纵效果良好。

单因素方差分析结果显示, 认同威胁组的职场物化($M = 3.18$, $SD = 0.59$, 95% CI [3.07, 3.30])显著高于对照组($M = 3.00$, $SD = 0.63$, 95% CI [2.88, 3.12]), $F(1, 212) = 4.66$, $p = 0.032$, $\eta_p^2 = 0.02$ 。

为了验证控制感是否在感知机器人认同威胁对职场物化的影响中起中介作用, 使用 Hayes (2013)提供的 SPSS 插件 PROCESS (Model 4), 以感知机器人认同威胁为自变量(对照组 = -1, 认同威胁组 = 1), 控制感为中介变量, 职场物化为因变量, 设定 Bootstrap 样本量为 5000, 采用偏差校正的方法, 选取 95%置信区间进行中介效应检验。数据结果显示, 控制感的中介效应值为 0.02, 95%的 Bootstrap 置信区间为[0.003, 0.051], 不包含 0, 表明中介作用显著; 并且在控制中介变量后, 感知机器人认同威胁对职场物化的直接效应为 0.07, 95%的 Bootstrap 置信区间为[-0.012, 0.151], 包含 0, 表明其直接效应不再显著, 控制感在感知机器人认同威胁对职场物化的影响中起完全中介作用。

以职场物化为因变量考察感知机器人认同威胁(对照组 = -1, 认同威胁组 = 1)与外部能动性的交互作用, 结果表明感知机器人认同威胁和外部能动性对职场物化存在显著的交互作用($b = -0.10$, $SE = 0.05$, $t = -1.98$, $p = 0.049$), 感知机器人威胁组的职场物化显著高于对照组($b = 0.09$, $SE = 0.04$, $t = 2.17$, $p = 0.031$), 外部能动性的高低对职场物化无显著影响($b = 0.02$, $SE = 0.05$, $t = 0.44$, $p = 0.661$), 模型的 $\Delta R^2 = 0.02$, $F(1, 210) = 3.92$, $p = 0.049$ 。交互作用如图 3 所示, 简单斜率分析结果表明, 在低外部能动性条件下, 感知机器人认同威胁对职场物化的影响显著($b = 0.18$, $SE = 0.06$, $t = 2.63$, $p = 0.004$); 而在高外部能动性条件下, 感知机器人认同威胁对职场物化的影响不显著($b = 0.01$, $SE = 0.06$, $t = 1.10$, $p = 0.920$)。

4.3 研究 3c: 特定结构信念的调节作用

4.3.1 被试

使用 G*Power 3.1 软件来确定本研究所需的样本量, 根据已有关于机器人劳动力相关研究(Jackson et al., 2020), 取中等效应量 $f = 0.2$, 显著性水平

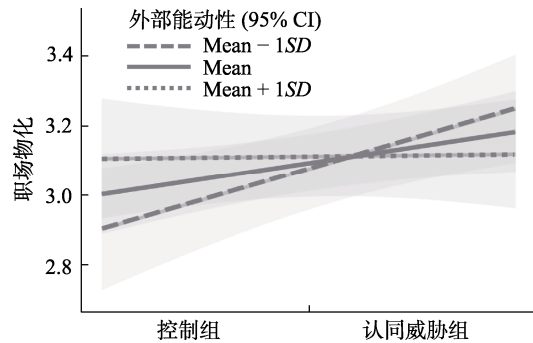


图 3 外部能动性的调节作用

$\alpha = 0.05$, 单因素两水平被试间设计需要共 200 名被试才能达到 80%的统计检验力。由于研究职场物化, 因此设置了填写条件为目前在职, 并且考虑到可能有被试无法通过注意检查, 并在 Credamo 平台上共计招募了 208 名被试。其中 8 名被试没有通过注意检查, 最终被试量为 200 人。被试平均年龄为 31.62 岁($SD = 7.44$), 女性 125 名(占 62.5%)。

4.3.2 过程

本研究为单因素两水平被试间设计, 被试被随机分配到认同威胁组和对照组, 在最终有效被试中, 认同威胁组 100 人, 对照组 100 人。认同威胁组和对照组的操纵过程同研究 2c。在完成写作任务后, 被试依次填写感知机器人威胁量表(同研究 1a, 在本实验中, 现实威胁维度的 Cronbach's α 系数为 0.81, 认同威胁维度的 Cronbach's α 系数为 0.89)、控制感问卷(同研究 2a, 在本实验中, 控制感全部条目的 Cronbach's α 系数为 0.84)、特定结构信念量表和职场物化量表(同研究 1a, 在本实验中该量表的 Cronbach's α 系数为 0.64)。

根据补偿控制理论中对肯定特定结构的说明, 如相信做出特定的行为能够可靠地产生预期结果, 并且这种清晰、可靠的“行为—结果”可能性是针对减少控制感的具体情境的, 结合本研究目的及具体情境, 编制特定结构信念测量条目如下: (1)我认为拔掉电源可以控制机器人; (2)我认为编写代码可以控制机器人; (3)我认为破坏机器人的系统可以控制机器人; (4)我认为有特定的行为可以控制机器人。以上测量条目均为 Likert 7 点量表, 1 为非常不同意, 7 为非常同意。本研究中该量表的 Cronbach's α 系数为 0.62。量表中穿插有注意检查题项如“本题请选择非常不同意”, 最后填写人口统计学信息。

4.3.3 结果

单因素方差分析结果显示, 认同威胁组的现实威胁评分($M = 4.89$, $SD = 0.86$, 95% CI [4.72, 5.06])

也显著高于对照组($M = 3.87$, $SD = 0.93$, 95% CI [3.69, 4.06]), $F(1, 198) = 63.85$, $p < 0.001$, $\eta_p^2 = 0.24$; 认同威胁组的认同威胁评分($M = 5.07$, $SD = 1.06$, 95% CI [4.86, 5.28])显著高于对照组($M = 3.80$, $SD = 1.21$, 95% CI [3.56, 4.04]), $F(1, 198) = 61.84$, $p < 0.001$, $\eta_p^2 = 0.24$ 。操纵效果良好。

单因素方差分析结果显示, 认同威胁组的职场物化($M = 3.32$, $SD = 0.65$, 95% CI [3.19, 3.44])显著高于对照组($M = 3.10$, $SD = 0.66$, 95% CI [2.97, 3.23]), $F(1, 198) = 5.59$, $p = 0.019$, $\eta_p^2 = 0.03$ 。

为了重复验证控制感是否在感知机器人认同威胁对职场物化的影响中起中介作用, 使用 Hayes (2013)提供的 SPSS 插件 PROCESS (Model 4), 以感知机器人认同威胁为自变量(对照组 = -1, 认同威胁组 = 1), 控制感为中介变量, 职场物化为因变量, 设定 Bootstrap 样本量为 5000, 采用偏差校正的方法, 选取 95%置信区间进行中介效应检验。数据结果显示, 控制感的中介效应值为 0.03, 95%的 Bootstrap 置信区间为[0.003, 0.067], 不包含 0, 表明中介作用显著; 并且在控制中介变量后, 感知机器人认同威胁对职场物化的直接效应为 0.08, 95%的 Bootstrap 置信区间为[-0.005, 0.169], 包含 0, 表明其直接效应不再显著, 控制感在感知机器人认同威胁对职场物化的影响中起完全中介作用。

以职场物化为因变量考察感知机器人认同威胁(对照组 = -1, 认同威胁组 = 1)与特定结构信念的交互作用, 结果表明感知机器人认同威胁和特定结构信念对职场物化存在显著的交互作用($b = -0.14$, $SE = 0.05$, $t = -2.76$, $p = 0.006$), 感知机器人威胁组的职场物化显著高于对照组($b = 0.11$, $SE = 0.05$, $t = 2.38$, $p = 0.018$), 特定结构信念的高低对职场物化无显著影响($b = 0.01$, $SE = 0.05$, $t = 0.25$, $p = 0.806$), 模型的 $\Delta R^2 = 0.04$, $F(1, 196) = 7.63$, $p = 0.006$ 。交互作用如图 4 所示, 简单斜率分析结果表明, 在低特定结构信念条件下, 感知机器人认同威胁对职场物化的影响显著($b = 0.24$, $SE = 0.07$, $t = 3.64$, $p < 0.001$); 而在高特定结构信念条件下, 感知机器人认同威胁对职场物化的影响不显著($b = -0.02$, $SE = 0.07$, $t = -0.27$, $p = 0.784$)。

5 讨论

通过 3 项研究(共 8 项子研究), 结果发现, 感知机器人威胁会增加职场物化, 控制感在其中起中介作用, 并且补偿控制的另外三种策略, 即加强个

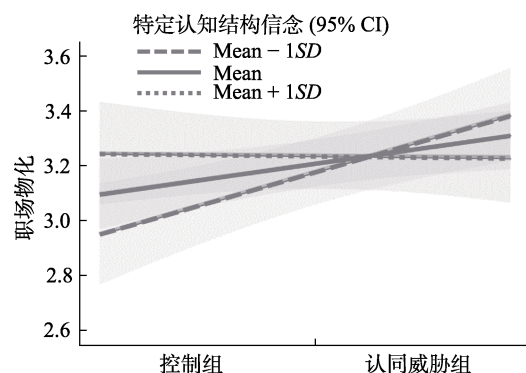


图4 特定结构信念的调节作用

人能动性、支持外部能动性以及肯定特定结构, 能够调节感知机器人威胁对职场物化的影响。具体而言, 首先通过对机器人公司员工的大规模问卷调查发现了感知机器人威胁与职场物化的正相关关系(研究 1a), 然后又通过多种研究方法反复验证感知机器人威胁会增加职场物化(研究 1b~3c)。通过对被试控制感的测量, 进一步发现了控制感是感知机器人威胁影响职场物化的中介机制, 即感知到机器人威胁降低了被试的控制感, 进而通过职场物化来进行补偿(研究 2a~2c), 并且发现主要是感知机器人认同威胁造成了职场物化的差异(研究 2c)。通过对其他三种补偿控制策略的测量, 研究还发现了个人能动性(研究 3a)、外部能动性(研究 3b)以及特定结构信念(研究 3c)对感知机器人威胁影响职场物化的调节作用。

更为有趣的是, 如果将感知机器人威胁细化, 本研究实际上发现了感知机器人认同威胁对职场物化的效应要胜于感知机器人现实威胁, 同时控制感中介的是认同威胁到职场物化间的路径。感知机器人认同威胁更多导致了职场物化增加的结果在一定程度上回应了前人对机器人威胁的相关研究结果的差异, 如 Huang 等人(2021)的研究发现相对于机器人现实威胁, 认同威胁更能预测人们对机器人的负面态度。这一结论也不是绝对的, 如 Zlotowski 等人(2017)则发现机器人现实威胁和认同威胁都是导致对机器人的负面态度和对机器人研究的反对增加的中介因素, 不过现实威胁的中介效应还更强。本研究实际上均发现了现实威胁和认同威胁对职场物化的效应, 但认同威胁效应更强, 且控制感只能中介认同威胁, 控制补偿的其他三种策略也只能调节认同威胁对职场物化的影响。这都表明, 机器人认同威胁是增加职场物化的主要因素, 也是导致控制补偿的主要机制。这可能与本研究所探讨的因变

量职场物化有关,由于“职场”这一概念本身就与对工作的现实威胁密切相关,因此认同威胁在一定程度上可能已然包含了现实威胁,且较其程度更甚。

而机器人认同威胁在职场物化情境下,除了越来越多地取代人类的工作,是因为模糊了人与机器之间的界限,从而造成对人类独特性的威胁(Yogeeswaran et al., 2016)。在同样的情境下,实际上人类独特性可以被威胁,也可以被凸显,这似乎是一件事情的两面,从何种角度看待则会造成不同的结果。如也有研究发现机器人或许能够提升人们的泛人类主义,更多地将人类同胞视为自己的内群体从而减少偏见,产生正面人际效应(Jackson et al., 2020)。本研究从感知机器人威胁的角度切入,发现其会增加职场物化,从而验证了机器人对人际关系可能的负面影响,这明显与 Jackson 等人(2020)所发现的机器人对人际关系的正面影响是相反的。本研究也对泛人类主义进行了测量,但并未发现与前人研究类似的积极效应(Jackson et al., 2020)。这一差异可能是由于职场的竞争特征,导致从“威胁”角度来看待同一问题的可能性要大于“凸显”角度。当然,本研究虽然实验操纵材料基本一致(本文所使用的实验操纵材料为中文翻译版),但在 Jackson 等人(2020)所探讨的自变量仅为机器人的“凸显”,并未涉及机器人“威胁”,但本文被试在阅读实验操纵材料后填写了感知机器人威胁量表作为操纵检查,这可能也在一定程度上对机器人威胁进行了提醒和强调。本研究的结果实际上与之前的机器人影响人类之间关系的研究结果相似,如研究发现,由于机器人对人类社会有限资源的挤占,机器人的增加在一定程度上会使人与人的关系更加紧张而非更加积极(如 Frey et al., 2018; Im et al., 2019)。此外,从相对意义而言,物化他人的同时也会造成自身与他人的区隔,即在物化他人的同时一定程度上激发了自身的优越感,产生“对比”效应。本文关注的重点在于对他人的物化,未来研究可以更加深入探讨人们在受到机器人威胁时产生的自身相对优势问题。

机器人认同威胁导致的这些结果也有一定的实践意义。机器人渗入职场对人们的工作、资源等造成的现实威胁或许很难避免,但机器人对于人之为人的独特性造成的认同威胁却能够在一定程度上减少。例如如今很多机器人都被设计为具有拟人化的特征,很多研究也发现机器人拟人化有积极影响,如拟人化的机器人更受人们信赖(Rehm & Krogsager, 2013)、能够得到更多表扬和更少惩罚

(Bartneck et al., 2007)等。然而,机器人更高层次的拟人化同时也会使人类知觉到更严重的威胁,尤其是其能力高于人类时(Yogeeswaran et al., 2016)。本研究的发现进一步探讨了人类感知到机器人威胁,尤其是感知到机器人的认同威胁可能会导致的人际关系负面后果,即职场物化。那么既然机器人过于像人会对人类造成认同威胁,且这种威胁会产生消极影响,那么在对机器人的设计中就必须把握好拟人化的尺度,以预防其消极后果的产生。

尽管本文通过一系列研究对感知机器人威胁影响职场物化的现象、心理机制和边界条件进行了探讨,但研究仍存在一定程度的局限性,有待在后续研究中进行改善和提升。首先职场物化的方向有待拓展,本文主要探讨的是平行职场物化,即同事之间的职场物化,但其实职场物化的方向还可能是下行(如领导对员工; Gruenfeld et al., 2008; Landau et al., 2012)和上行(如员工对领导; Inesi et al., 2014),不同的物化方向可能会带来控制感的差异,从而影响职场物化的结果。其次,研究的生态效度仍有提升空间,尽管本文的 8 个子研究尽量增加了被试群体来源的多样性,如不仅涵盖大学生被试,还有从网络平台上招募的来自全国各地的被试,并且也有真正在机器人工厂工作的一线工作者,但是之后的研究可以更多地涵盖不同行业的可能接触到机器人的工作者,对本文的研究结果进行重复验证。此外本文主要采用了问卷研究和网络实验,并且对职场物化的测量也采用了自我报告的方式,没有对被试的真实行为进行观察记录,未来可以通过现场实验等方法,让被试与真实的机器人现场互动从而产生感知威胁,在更加真实的社会情境中观察被试的真实行为。在实验设计的细节方面,文中 2.3 等研究中所使用的写作任务对照组的指导语在情绪属性上较为积极,这可能会对研究结果产生一定影响,在未来研究中应尽量避免,代之以更加中性的指导语。

6 结论

本研究结论如下:第一,感知机器人威胁会增加职场物化,并且感知机器人认同威胁的影响更强;第二,控制感在感知机器人威胁(主要是认同威胁)影响职场物化中起中介作用;第三,补偿控制的另外三种策略,即加强个人能动性、支持外部能动性以及肯定特定结构,能够调节感知机器人威胁对职场物化的影响。

参 考 文 献

- Adler, N. E., Epel, E. S., Castellazzo, G., & Ickovics, J. R. (2000). Relationship of subjective and objective social status with psychological and physiological functioning: Preliminary data in healthy white women. *Health Psychology, 19*(6), 586–592.
- Baldissarri, C., Valtorta, R. R., Andrighetto, L., & Volpato, C. (2017). Workers as objects: The nature of working objectification and the role of perceived alienation. *Testing, Psychometrics, Methodology in Applied Psychology, 24*, 153–166.
- Bartneck, C., Reichenbach, J., & Carpenter, J. (2007). Use of praise and punishment in human-robot collaborative teams. *ROMAN 2006 - The 15th IEEE International Symposium on Robot and Human Interactive Communication* (pp.177–182). Hatfield, UK.
- Beck, J. T., Rahinel, R., & Bleier, A. (2020). Company worth keeping: Personal control and preferences for brand leaders. *Journal of Consumer Research, 46*(5), 871–886.
- Belmi, P., & Schroeder, J. (2021). Human “resources”? Objectification at work. *Journal of Personality and Social Psychology, 120*(2), 384–417.
- Borenstein, J. (2011). Robots and the changing workforce. *AI & Society, 26*(1), 87–93.
- Calogero, R. M. (2013). Objects don’t object: Evidence that self-objectification disrupts women’s social activism. *Psychological Science, 24*(3), 312–318.
- 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.
- Frey, C. B., Berger, T., & Chen, C. (2018). Political machinery: Did robots swing the 2016 U.S. presidential election? *Oxford Review of Economic Policy, 34*(3), 418–442.
- Gervais, S. J., DiLillo, D., & McChargue, D. (2014). Understanding the link between men’s alcohol use and sexual violence perpetration: The mediating role of sexual objectification. *Psychology of Violence, 4*(2), 156–169.
- Greenaway, K. H., Louis, W. R., Hornsey, M. J., & Jones, J. M. (2014). Perceived control qualifies the effects of threat on prejudice. *British Journal of Social Psychology, 53*(3), 422–442.
- Gruenfeld, D. H., Inesi, M. E., Magee, J. C., & Galinsky, A. D. (2008). Power and the objectification of social targets. *Journal of Personality and Social Psychology, 95*(1), 111–127.
- Hayes, A. F. (2013). *Introduction to mediation, moderation, and conditional process analysis: A regression-based approach*. New York: Guilford Press.
- Huang, H. L., Cheng, L. K., Sun, P. C., & Chou, S. J. (2021). The effects of perceived identity threat and realistic threat on the negative attitudes and usage intentions toward hotel service robots: The moderating effect of the robot’s anthropomorphism. *International Journal of Social Robotics, 13*, 1599–1611.
- Im, Z. J., Mayer, N., Palier, B., & Rovny, J. (2019). The “losers of automation”: A reservoir of votes for the radical right? *Research & Politics, 6*, 1–7.
- Inesi, M. E., Lee, S. Y., & Rios, K. (2014). Objects of desire: Subordinate ingratiation triggers self-objectification among powerful. *Journal of Experimental Social Psychology, 53*, 19–30.
- Jackson, J. C., Castelo, N., & Gray, K. (2020). Could a rising robot workforce make humans less prejudiced? *American Psychologist, 75*(7), 969–982.
- Kant, I. (1996). *The metaphysics of morals* (M. Gregor, Trans.). Cambridge, UK: Cambridge University Press. (Original work published 1797)
- Kay, A. C., Gaucher, D., Napier, J. L., Callan, M. J., & Laurin, K. (2008). God and the government: Testing a compensatory control mechanism for the support of external systems. *Journal of Personality and Social Psychology, 95*(1), 18–35.
- Kay, A. C., Whitson, J. A., Gaucher, D., & Galinsky, A. D. (2009). Compensatory control: Achieving order through the mind, our institutions, and the heavens. *Current Directions in Psychological Science, 18*(5), 264–268.
- Lachman, M. E., & Weaver, S. L. (1998). The sense of control as a moderator of social class differences in health and well-being. *Journal of Personality and Social Psychology, 74*(3), 763–773.
- Landau, M. J., Kay, A. C., & Whitson, J. A. (2015). Compensatory control and the appeal of a structured world. *Psychological Bulletin, 141*(3), 694–722.
- Landau, M. J., Sullivan, D., Keefer, L. A., Rothschild, Z. K., & Osman, M. R. (2012). Subjectivity uncertainty theory of objectification: Compensating for uncertainty about how to positively relate to others by downplaying their subjective attributes. *Journal of Experimental Social Psychology, 48*(6), 1234–1246.
- Leo, X., & Huh, Y. E. (2020). Who gets the blame for service failures? Attribution of responsibility toward robot versus human service providers and service firms. *Computers in Human Behavior, 113*(4), 106520.
- Marx, K. (1964). *Early writings* (T. B. Bottomore, Trans.). New York, NY: McGraw-Hill. (Original work published 1844)
- McClure, P. K. (2018). “You’re fired,” says the robot: The rise of automation in the workplace, technophobes, and fears of unemployment. *Social Science Computer Review, 36*(2), 139–156.
- McFarland, S., Webb, M., & Brown, D. (2012). All humanity is my ingroup: A measure and studies of identification with all humanity. *Journal of Personality and Social Psychology, 103*(5), 830–853.
- Murashov, V., Hearl, F., & Howard, J. (2016). Working safely with robot workers: Recommendations for the new workplace. *Journal of Occupational and Environmental Hygiene, 13*(3), D61–D71.
- Nussbaum, M. C. (1995). Objectification. *Philosophy & Public Affairs, 24*(4), 249–291.
- Nussbaum, M. C. (1999). *Sex and social justice*. New York, NY: Oxford University Press.
- Orehek, E., & Weaverling, C. G. (2017). On the nature of objectification: Implications of considering people as means to goals. *Perspectives on Psychological Science, 12*(5), 719–730.
- Rehm, M., & Krogsager, A. (2013). Negative affect in human robot interaction—Impoliteness in unexpected encounters with robots. In *Proceedings of the 2013 IEEE international workshop on robot and human communication* (pp. 45–50). Gyeongju, Korea.
- Saguy, T., Quinn, D. M., Dovidio, J. F., & Pratto, F. (2010). Interacting like a body: Objectification can lead women to narrow their presence in social interactions. *Psychological Science, 21*(2), 178–182.
- Schmidt, A. F., & Kistemaker, L. M. (2015). The sexualized-body-inversion hypothesis revisited: Valid indicator of sexual objectification or methodological artifact? *Cognition, 134*, 77–84.
- Shepherd, S., & Kay, A. C. (2018). Guns as a source of order

- and chaos: Compensatory control and the psychological (dis)utility of guns for liberals and conservatives. *Journal of the Association for Consumer Research*, 3(1), 16–26.
- Skowronski, M., Busching, R., & Krahé, B. (2021). The effects of sexualized video game characters and character personalization on women's self-objectification and body satisfaction. *Journal of Experimental Social Psychology*, 92(1), 104051.
- Smith, A., & Anderson, M. (2017). *Automation in everyday life*. Washington, DC: Pew Research Center.
- Stephan, W. G., Ybarra, O., & Rios, K. (2016). Intergroup threat theory. In T. D. Nelson (Ed.), *Handbook of prejudice, stereotyping, and discrimination* (2nd ed., pp. 255–278). Mahwah, NJ: Lawrence Erlbaum Associates.
- Taskin, L., Parmentier, M., & Stinglhamber, F. (2019). The dark side of office designs: Towards de-humanization. *New Technology, Work and Employment*, 34(3), 262–284.
- Thompson, S. C., & Schlehofer, M. M. (2008). Control, denial, and heightened sensitivity reactions to personal threat: Testing the generalizability of the threat orientation approach. *Personality and Social Psychology Bulletin*, 34(8), 1070–1083.
- 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., Yu, F., Wu, J., Han, T., & Zhao, L. (2017). Anthropomorphism: Antecedents and consequences. *Advances in Psychological Science*, 25(11), 1942–1954.
- [许丽颖, 喻丰, 邬家骅, 韩婷婷, 赵靛. (2017). 拟人化: 从“它”到“他”. *心理科学进展*, 25(11), 1942–1954.]
- Yang, S., Guo, Y., Hu, X., Shu, S., & Li, J. (2016). Do lower class individuals possess higher levels of system justification? An examination from the social cognitive perspectives. *Acta Psychologica Sinica*, 48(11), 1467–1478.
- [杨沈龙, 郭永玉, 胡小勇, 舒首立, 李静. (2016). 低阶层者的系统合理化水平更高吗?——基于社会认知视角的考察. *心理学报*, 48(11), 1467–1478.]
- Yogeeswaran, K., Zlotowski, J., Livingstone, M., Bartneck, C., Sumioka, H., & Ishiguro, H. (2016). The interactive effects of robot anthropomorphism and robot ability on perceived threat and support for robotics research. *Journal of Human-Robot Interaction*, 5(2), 29–47.
- Yu, F. (2020). On AI and human beings. *Frontiers*, (1), 30–36.
- [喻丰. (2020). 人工智能与人之为人. *学术前沿*, (1), 30–36.]
- Yu, F., & Xu, L. (2018). How to make an ethical artificial intelligence? Answer from a psychological perspective. *Global Media Journal*, 5(4), 24–42.
- [喻丰, 许丽颖. (2018). 如何做出道德的人工智能体? 心理学的视角. *全球传媒学刊*, 5(4), 24–42.]
- Yu, F., & Xu, L. (2020). Artificial intelligence and anthropomorphism. *Journal of Northwest Normal University (Social Sciences)*, 57(5), 52–60.
- [喻丰, 许丽颖. (2020). 人工智能之拟人化. *西北师大学报(社会科学版)*, 57(5), 52–60.]
- Zhang, Y., Xu, L., Yu, F., Ding, X., Wu, J., & Zhao, L. (2022). A three-dimensional motivation model of algorithm aversion. *Advances in Psychological Science*, 30(5), 1093–1105.
- [张语嫣, 许丽颖, 喻丰, 丁晓军, 邬家骅, 赵靛. (2022). 算法拒绝的三维动机理论. *心理科学进展*, 30(5), 1093–1105.]
- Zhou, H., & Long, L. (2004). Statistical remedies for common method biases. *Advances in Psychological Science*, 12(6), 942–950.
- [周浩, 龙立荣. (2004). 共同方法偏差的统计检验与控制方法. *心理科学进展*, 12(6), 942–950.]
- Zlotowski, J., Yogeeswaran, K., & Bartneck, C. (2017). Can we control it? Autonomous robots threaten human identity, uniqueness, safety, and resources. *International Journal of Human-Computer Studies*, 100, 48–54.

The influence of perceived robot threat on workplace objectification

XU Liying¹, WANG Xuehui², YU Feng¹, PENG Kaiping²

(¹ Department of Psychology, Wuhan University, Wuhan 430072, China)

(² Department of Psychology, School of Social Sciences, Tsinghua University, Beijing 100084, China)

Abstract

With buzzwords such as “tool man”, “laborer” and “corporate slave” sweeping the workplace, workplace objectification has become an urgent topic to be discussed. With the increasing use of artificial intelligence, especially robots in the workplace, the workplace effects produced by robots are also worth paying attention to. Therefore, the present paper aims to explore whether people's perception of robots' threat to them will produce or aggravate workplace objectification. On the basis of reviewing the related research on workplace objectification and robot workforce, and combined with intergroup threat theory, this paper elaborates the realistic threat to human employment and security caused by robot workforce, as well as the identity threat to human identity and uniqueness. From the perspective of compensatory control theory, this paper proposes the deep mechanisms and boundary conditions of that perceiving robot threat will reduce people's sense of control, thereby stimulating the control compensation mechanism, which in turn leads to workplace objectification.

This research is composed of eight studies. The first study includes two sub-studies, which investigate the

relationship between perceived robot threat and workplace objectification through questionnaires and online experiments. This study tries to find a positive correlation and a causal association between perceived robot threat and workplace objectification. The second study includes three sub-studies, which explore why perceived robot threat increases workplace objectification. This study tries to verify the mediating effect of control compensation (sense of control), to explain the psychological mechanism behind the effect of perceived robot threat on workplace objectification, and to repeatedly verify it through different research methods. The third study includes three sub-studies. Based on the three compensatory control strategies proposed by the control compensation theory in addition to affirming nonspecific structure, this study tries to further explore the moderating effect of personal agency, external agency, and specific structure.

The main findings of this paper are as follows. First, perceived robot threat will increase workplace objectification, and perceived robot identity threat has a stronger effect. Second, the sense of control plays a mediating role in the effect of perceived robot threat (mainly identity threat) on workplace objectification. Specifically, the higher the perceived robot identity threat, the lower the sense of control, and the more serious the workplace objectification. Third, the other three strategies proposed by compensatory control theory, namely strengthening personal agency, supporting external agency and affirming specific structure, can moderate the effect of perceived robot threat on workplace objectification.

The main theoretical contributions of this paper are as follows. First, it reveals the negative influence of robots on interpersonal relationships and their psychological mechanism. Second, it extends the applicability of compensatory control theory to the field of artificial intelligence, proposing and verifying that perceived robot threat increases workplace objectification through compensatory control. Third, the relationship between different compensation control strategies is discussed, and the moderating model of perceived robot threat affecting workplace objectification is proposed and verified. The main practical contributions are: first, it provides insights into the anthropomorphic design of robots; second, it helps to better understand, anticipate and mitigate the negative social impact of robots.

Keywords workplace objectification, perceived robot threat, realistic threat, identity threat, compensatory control theory