

中文引用格式:牛莉霞,李波,李国. AI焦虑的安全管理困境:人机信任与系统透明度机制[J]. 中国安全科学学报, 2026, 36(3): 33-40.

英文引用格式:NIU Lixia, LI Bo, LI Guo. Safety management dilemma under AI anxiety: human-AI trust and system transparency mechanism[J]. China Safety Science Journal, 2026, 36(3): 33-40.

AI 焦虑的安全管理困境: 人机信任与系统透明度机制*

牛莉霞¹教授, 李波¹, 李国²

(1 辽宁工程技术大学 工商管理学院, 辽宁 葫芦岛 125105;

2 中矿检测(辽宁)有限公司, 辽宁 阜新 123000)

中图分类号: X912.9

文献标志码: A

DOI: 10.16265/j.cnki.issn1003-3033.2026.03.0299

基金项目: 国家自然科学基金资助(52174184); 辽宁省教育厅基本科研项目(LJ212510147042)。

【摘要】 为解决人工智能(AI)在企业安全管理与决策实践中因系统复杂性提升、信息不确定性增强而诱发的员工AI焦虑问题, 本文基于不确定管理理论(UMT), 构建“AI焦虑-人机信任-人机协同决策质量”的机制模型, 并引入系统透明度作为边界条件, 解释在风险提示与算法黑箱等情境下员工对AI的威胁评估与应对方式; 同时以523名AI应用企业员工为样本开展问卷调查, 采用验证性因子分析(CFA)与结构方程模型对测量模型、路径关系及调节效应进行检验, 并控制性别、年龄、教育、岗位性质与AI使用频率等变量。结果显示: AI焦虑降低人机协同决策质量; 人机信任在AI焦虑与协同决策质量之间发挥部分中介作用; 系统透明度正向调节人机信任对协同决策质量的影响, 高透明度可促进信任向高质量协同转化。

【关键词】 人工智能(AI)焦虑; 安全管理; 人机信任; 系统透明度; 人机协同决策质量; 不确定管理理论(UMT)

Safety management dilemma under AI anxiety: human-AI trust and system transparency mechanism

NIU Lixia¹, LI Bo¹, LI Guo²

(1 School of Business Administration, Liaoning Technical University, Huludao Liaoning 125105, China;

2 Zhongkuang Testing (Liaoning) Co., Ltd., Fuxin Liaoning 123000, China)

Abstract: To address employee' AI anxiety arising from increased system complexity and heightened information uncertainty in the application of AI to corporate safety management and decision-making, a mechanism model of "AI anxiety-human-AI trust-human-AI collaborative decision quality" was developed based on UMT. It introduced system transparency as a boundary condition to explain how employees appraised and coped with AI-related threats in contexts such as risk warnings and algorithmic black boxes. A questionnaire survey was conducted with a sample of 523 employees from AI-adopting enterprises. Confirmatory factor analysis (CFA) and structural equation modeling were employed to test the measurement model, path relationships, and moderating effects, while controlling for variables such as gender, age, education, job type, and AI usage frequency. The results show that AI anxiety reduces the

quality of human-AI collaborative decision-making. Human-AI trust partially mediates the relationship between AI anxiety and collaborative decision-making quality. System transparency positively moderates the effect of human-AI trust on collaborative decision-making quality, such that higher transparency facilitates the translation of trust into higher-quality collaboration.

Keywords: artificial intelligence (AI) anxiety; safety management; human-AI trust; system transparency; human-AI collaborative decision quality; uncertainty management theory (UMT)

0 引言

在新一轮科技革命与产业变革推动下,人工智能(Artificial Intelligence, AI)深度嵌入高风险、高复杂性组织环境,成为企业安全管理数智化转型的关键动力。《“十四五”国家科技创新规划》将 AI 定位为国家战略科技力量突破口,强调其在风险分级管控、隐患治理、作业许可与联锁、智能预警与应急响应等环节的应用^[1]。AI 驱动的识别—评估—预警—处置—复盘闭环,推动安全管理由经验驱动向数据驱动、智能协同转变^[2-3]。但 AI 系统的复杂性、不可解释性与结果不确定性也易引发员工对可靠性与可控性的担忧,即 AI 焦虑,并表现为不信任预警、绕开联锁、规避或延迟采纳建议等,进而削弱人机协同决策质量与安全治理闭环成效^[4-6]。因此,从安全科学与安全管理视角揭示 AI 焦虑在高风险场景中的作用机制与边界条件,具有重要理论意义与实践价值。

近年来,学术界对 AI 焦虑的研究持续深化。JOHNSON 等^[4]厘清其概念根源,指向对算法自主性、系统不透明性及人类控制力弱化的感知;LI Jian 等^[6]基于整合性恐惧获得理论提出多维结构与测量框架,强调风险感知、失控感与不确定线索的作用路径;SUSENO 等^[5]在人力资源样本中揭示了信念—焦虑—变革就绪度链条,并指出高绩效工作体系具有缓冲效应。郭娟等^[7]发现 AI 焦虑会降低员工与 AI 系统的互动频率、削弱协作潜力;刘小禹等^[8]指出情绪调节能力在 AI 焦虑与工作适应之间发挥关键中介作用;杜佩佳等^[9]揭示了 AI 技术对员工心理层面的双重影响;蒋建武等^[10]通过元分析证实 AI 焦虑与员工工作适应能力显著负相关。尽管如此,现有研究多聚焦心理适应与技术采纳,尚缺乏对其在高风险安全决策情境中作用机制的系统探讨,如在多主体协同决策过程中, AI 焦虑可能通过削弱员工对 AI 系统的信任,影响信息采纳、判断参与与协作意愿,从而降低人机协同决策质

量^[11]。此外,系统透明度作为可设计的技术属性虽被认为有助于信任形成^[2],但其是否以及如何调节人机信任—人机协同决策质量的关系仍缺乏明确结论。

鉴于此,笔者拟基于不确定管理理论(Uncertainty Management Theory, UMT),聚焦安全管理情境下的人机协同决策过程,构建 AI 焦虑—人机信任—人机协同决策质量的核心分析链条,并进一步检验系统透明度的调节作用,以期深化 AI 焦虑在安全管理领域的机制解释,并为企业优化人机协同决策流程、提升智能化安全治理能力提供理论依据。

1 AI 焦虑的安全管理作用机制与研究假设

1.1 AI 焦虑与人机协同决策质量

在人机协同决策背景下,决策质量是指人与 AI 系统协作过程中形成决策的准确性、可靠性与有效性^[11-12]。基于 UMT, AI 焦虑会加剧个体对决策风险的不确定感,强化面对 AI 时的警惕与防御心理^[13-14],使其更倾向于质疑系统输出、保留甚至忽视 AI 信息,回避采纳建议并转而依赖经验或直觉判断。由此产生的认知与行为偏差不仅降低人机协同决策质量,还可能导致关键安全信息被忽略或误判,进而增加安全事故风险^[15-16]。在安全管理领域,高质量协同决策是降低事故风险、提升管理效率的关键^[17];在人机融合实践中,其既影响事故预防与风险控制成效,也关系着安全治理体系能否高效运转。因此,揭示人机协同决策的作用机制与优化路径,是安全管理智能化、精细化发展的内在要求,也是提升高风险行业安全水平的重要突破口。

因此,提出以下假设:

H₁: AI 焦虑负向影响人机协同决策质量。

1.2 人机信任的中介作用

人机信任是指个体在与 AI 系统交互过程中对

系统可靠性、公正性及功能性的信任。信任一直是管理学和心理学领域的重要研究主题,经典信任理论由 MAYER 等^[18]提出。在人机协同决策过程中,人机信任是促进信息共享和深度协作的关键心理机制,对提升决策质量具有重要作用。根据 UMT, AI 焦虑源于个体对 AI 技术的不确定感,这种不确定感会增强个体的戒备心理^[13-14],使其对 AI 系统的可靠性和公正性产生质疑,并减少对 AI 决策的依赖,从而降低对 AI 系统的信任^[19]。因此,提出以下假设:

H₂: AI 焦虑与人机信任之间存在负相关关系。

一方面,人机信任下降会削弱协作深度与互动频率,并降低个体对 AI 信息的采纳意愿,从而削弱人机协同决策效果;由此引发的决策延误与判断偏差,可能直接威胁企业运行的安全稳定。相反,人机信任可促进信息共享与深度协作,提升决策的准确性、可靠性与安全性^[20-21];个体对 AI 系统越信任,越倾向于采纳系统建议并积极参与协作^[22]。基于此,进一步提出以下假设:

H₃: 人机信任与人机协同决策质量之间存在正相关关系。

结合 H₂ 与 H₃ 可知: AI 焦虑可能通过降低员工对 AI 系统的信任,进而间接削弱人机协同决策质量。因此提出:

H₄: 人机信任在人机协同决策质量与 AI 焦虑之间起中介作用。

1.3 系统透明度的调节作用

系统透明度是指用户对 AI 系统决策逻辑、运作方式及潜在风险点的理解程度^[23]。在安全敏感决策环境中,透明度水平不仅可能改变人机信任对协同决策质量的作用强度,也影响员工对安全指令与预警信号的响应速度和准确性。基于 UMT,系统透明度可通过降低不确定性、改善信息交流,促进信任基础上的人机合作,从而提升协同决策质量^[24],对复杂技术系统的安全运行与风险防控具有重要意义。

具体而言,当系统透明度较高时, AI 的决策逻辑、操作流程与风险提示更清晰可理解,员工更易判断系统的安全性与可靠性,减少不确定感与焦虑^[25],从而强化人机信任对协同决策质量的正向作用,促使员工更愿意采纳建议并开展更充分的人机互动,进而提升协同决策质量^[26]。相反,当系统透明度较低时,员工难以及时把握决策依据与潜在风险信息,易产生认知模糊与理解偏差,安全感知下降

并加剧焦虑与不确定感;此时人机信任更难建立与维持,即便存在一定信任,也难以有效转化为高质量的决策互动与协作,最终削弱人机协同决策效果。因此,提出以下假设:

H₅: 系统透明度在人机信任与人机协同决策质量之间起正向调节作用。

基于上述分析,构建以 AI 焦虑为自变量、人机信任为中介变量、人机协同决策质量为因变量,并引入系统透明度作为调节变量的理论模型(图 1)。

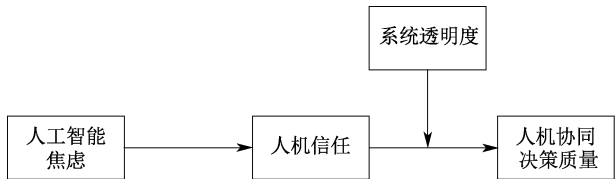


图 1 理论模型
Fig. 1 Theoretical model

2 人机信任与系统透明度实证研究

2.1 人机信任与系统透明度样本分析

正式调研前,组织约 50 名来自能源、制造及高科技行业的员工开展预试验,检验条目语义清晰度与情境适用性并据反馈优化表述。随后进行信度与效度检验:4 个核心量表 Cronbach's α 均大于 0.6、复合信度均大于 0.7;验证性因子分析(Confirmatory Factor Analysis, CFA)结果显示,各构念平均方差提取量(Average Variance Extracted, AVE)均大于 0.5,且 AVE 平方根均大于变量间相关系数,表明量表具有良好的信效度,为后续实证分析奠定基础。

正式调研采用多阶段问卷调查,面向全国 AI 技术应用企业员工收集数据,调查前说明仅用于科研并严格匿名处理;为提高回收率与 3 阶段数据匹配完整性,设置唯一抽奖编号用于匹配 3 次问卷,完成 3 次调查者给予手机充值卡或管理学书籍等奖励。

为降低同源偏差影响,采用 3 阶段纵向设计,间隔均为 1 个月收集数据:T1 测量员工 AI 焦虑(发放 647 份,回收 621 份,回收率 95.98%);T2 面向 T1 回收样本测量人机信任与系统透明度感知(发放 621 份,回收 589 份,回收率 94.85%);T3 测量人机协同决策质量并收集人口统计信息(发放 589 份,回收有效问卷 523 份,有效回收率 88.79%)。最终获得 523 份 3 阶段匹配样本:男性 51.6%、女性 48.4%;18~55 岁年龄段中,26~35 岁的占 55.1%;本科及以上学历 67.5%。样本覆盖制造业、高科技、医疗保健、金融、能源与交通运输等行业。

2.2 AI 焦虑与人机信任、系统透明度量表测量

量表均选取自国内外权威期刊中已验证的成熟量表,并结合研究背景对部分题项进行情境化修订。针对国外量表,采用翻译—回译程序以保证语义一致性。所有量表均采用李克特 5 点计分(1=“完全不符合”,5=“完全符合”)。研究主要变量及其测量工具如下:

- 1) AI 焦虑。测量量表来源于 LI Jian 等^[6]的研究,共 6 道题项,典型题项如“我担心 AI 在工作中给我带来安全风险”。量表 Cronbach’s $\alpha=0.845$,表明内部一致性较高。
- 2) 人机信任。测量量表来源于 GULATI 等^[26]的研究,共 12 个题项,典型题项如“我相信 AI 将为我的最大利益行事”。量表 Cronbach’s $\alpha=0.868$,信度良好。
- 3) 人机协同决策质量。测量量表来源于 SHAMIM 等^[27]的研究,共 6 个题项,典型题项如“我认为我与 AI 一起作出了好的工作决策”和“我与 AI 一起的工作决策有助于提高工作绩效”。量表 Cronbach’s $\alpha=0.781$,具有较高可靠性。
- 4) 系统透明度。测量量表来源于 LI Jian 等^[6]的研究,共 3 个题项,典型题项如“我很清楚 AI 是如何作出决策的”。量表 Cronbach’s $\alpha=0.937$,内部一致性极高。

5) 控制变量。为排除人口统计因素干扰,将年龄、性别、教育程度、行业和工作年限作为控制变量纳入分析。相关研究表明:上述特征会影响技术接受、AI 焦虑及协同决策表现^[28-31],纳入控制变量有助于更准确识别核心变量间作用机制。

3 AI 焦虑与人机信任、系统透明度实证检验与结果

3.1 区分效度与共同方法偏差检验

为检验 AI 焦虑、人机信任、系统透明度与人机协同决策质量的区分效度,采用结构方程模型,并通过 AMOS 软件进行 CFA。结果(表 1)显示,4 因子模型拟合最佳卡方统计量 $\chi^2=337.324$,自由度 $df=318$, $\chi^2/df=1.061$,比较拟合指数(Comparative Fit Index, CFI)=0.997,塔克-刘易斯指数(Tucker-Lewis Index, TLI)=0.997,近似均方根误差(Root Mean Square Error of Approximation, RMSEA)=0.011,标准化均方根残差(Standardized Root Mean Square Residual, SRMR)=0.026,显著优于其他竞争模型,并通过 0.01 水平卡方检验,表明 4 个核心变量具有良好区分效度。此外,采用 Harman 单因素检验共同方法偏差,第一主成解释率为 37.895%,低于 40%临界值,说明不存在严重共同方法偏差,可开展后续分析。

表 1 CFA
Table 1 CFA

模型	χ^2	df	χ^2/df	CFI	TLI	RMSEA	SRMR
4 因子(J,X,T,C)	337.324	318	1.061	0.997	0.997	0.011	0.025 6
3 因子(J,X+T,C)	1 669.131	321	5.200	0.806	0.788	0.090	0.076 3
2 因子(J,X+T+C)	1 672.382	323	5.178	0.806	0.789	0.089	0.076 2
1 因子(J+X+T+C)	1 673.218	324	5.164	0.806	0.789	0.089	0.076 1

注:J 为 AI 焦虑,X 为人机信任,T 为系统透明度,C 为人机协同决策质量。

3.2 样本特征与相关性检验

各变量的均值、标准差及相关系数见表 2。相关分析结果显示,AI 焦虑与人机协同决策质量显著负相关(皮尔逊相关系数 $r=-0.607$,显著性概率 $p<$

0.01);AI 焦虑与人机信任显著负相关($r=-0.658$, $p<0.01$);人机信任与人机协同决策质量显著正相关($r=0.609$, $p<0.01$)。上述结果与理论预期一致,初步验证了研究假设。

表 2 均值、标准差和相关系数

Table 2 Mean, standard deviation and correlation coefficients

变量	均值	标准差	相关系数								
			1	2	3	4	5	6	7	8	9
性别	1.48	0.50	1	—	—	—	—	—	—	—	—
年龄	2.5	0.90	-0.016	1	—	—	—	—	—	—	—
工龄	2.48	1.23	-0.056	0.751**	1	—	—	—	—	—	—
学历	1.77	0.61	-0.025	0.007	0.038	1	—	—	—	—	—
行业	3.55	1.82	-0.019	-0.006	0.001	-0.045	1	—	—	—	—

续表 2

变量	均值	标准差	相关系数								
			1	2	3	4	5	6	7	8	9
AI 焦虑	2.45	0.91	0.057	0.114 **	0.149 **	0.039	-0.008	1	—	—	—
人机信任	3.55	0.78	-0.069	-0.147 **	-0.173 **	-0.040	-0.005	-0.658 **	1	—	—
系统透明度	2.43	1.08	0.069	0.120 **	0.080	0.062	0.038	0.160 **	-0.174 **	1	—
人机协同决策质量	3.56	0.84	-0.094 *	-0.129 **	-0.156 **	-0.062	-0.050	-0.607 **	0.609 **	-0.133 **	1

注: * 表示 $p < 0.05$, ** 表示 $p < 0.01$, *** 表示 $p < 0.001$, 表 3、表 5 同。

3.3 假设检验分析

3.3.1 主效应检验

在控制性别、年龄、工作年限、教育程度和所在行业的基础上,回归分析结果(表 3)显示, AI 焦虑

对 人 机 协 同 决 策 质 量 具 有 显 著 负 向 影 响 ($M4, \beta = -0.702, p < 0.001$; M1—M5 代表模型 1—模型 5), 验证了假设 H_1 。

表 3 人机信任的中介效应回归分析

Table 3 Regression analysis of the mediating effect of human-AI trust

变量类型	变量	人机信任		人机协同决策质量		
		模型 1	模型 2	模型 3	模型 4	模型 5
控制变量	性别	-0.079	-0.023	-0.104 *	-0.052 *	-0.041 *
	年龄	-0.036	-0.034	-0.024	-0.022	-0.006
	工龄	-0.149 *	-0.021	-0.141 *	-0.021	-0.011
	学历	-0.036	-0.007	-0.062	-0.034	-0.031
	行业	-0.008	-0.013	-0.055	-0.059	-0.053 *
自变量	AI 焦虑	—	-0.852 ***	—	-0.702 ***	-0.393 ***
中介变量	人机信任	—	—	—	—	0.480 ***
	决定系数 R^2	0.038	0.744	0.041	0.667	0.726
	解释力增量 ΔR^2	0.038	0.706	0.041	0.625	0.059
	统计量 F	4.081 ***	249.559 ***	4.467 ***	172.165 ***	194.932 ***
	F 变化量 ΔF	4.081 ***	1 420.913 ***	4.467 ***	968.842 ***	111.107 ***

3.3.2 中介效应检验

为检验人机信任的中介作用,采用因果步骤法与 Bootstrap 方法进行双重验证:

1) 因果步骤法。由表 3 可知: AI 焦虑显著负向影响人机信任 ($M2, \beta = -0.852, p < 0.001$), 验证了 H_2 ; 人机信任显著正向影响人机协同决策质量 ($M5, \beta = 0.480, p < 0.001$), 验证了 H_3 ; 控制 AI 焦虑后, 人机信任的正向影响显著 ($M5, \beta = 0.480, p < 0.001$), 同时, AI 焦虑对人机协同决策质量的负向影响依然显著 ($M5, \beta = -0.393, p < 0.001$), 说明人机信任在两者之间发挥部分中介作用, 验证了 H_4 。

2) Bootstrap 分析。AI 焦虑对人机协同决策质量 的 中 介 效 应 通 过 Bootstrap 法 进 一 步 检 验, 重 复 抽 样 5 000 次, 具 体 效 应 值 与 95% 置 信 区 间 见 表 4。结 果 显 示: AI 焦 虑 对 人 机 协 同 决 策 质 量 的 总 效 应 为 $-0.739\ 2$ ($95\% \text{CI} = [-0.785\ 9, -0.692\ 6]$); 直 接 效 应 为 $-0.362\ 1$ ($95\% \text{CI} = [-0.444\ 2, -0.280\ 0]$); 间 接 效 应 为 $-0.337\ 1$ ($95\% \text{CI} = [-0.451\ 2,$

$-0.307\ 9]$); 其 中, CI 表 示 置 信 区 间 (Confidence Interval)。总效应、直接效应和间接效应的置信区间均不包含 0, 进一步验证了人机信任的部分中介作用, H_4 得到稳健支持。

表 4 直接效应、间接效应和总效应分解

Table 4 Decomposition of direct, indirect and total effects

效应	效应值	标准误差	95%置信区间
间接效应	-0.337 1	0.035 7	$[-0.451\ 2, -0.307\ 9]$
直接效应	-0.362 1	0.041 8	$[-0.444\ 2, -0.280\ 0]$
总效应	-0.739 2	0.023 7	$[-0.785\ 9, -0.692\ 6]$

3.3.3 调节效应检验

为检验系统透明度在中介路径中的调节作用, 采用均值中心化处理后的人机信任与系统透明度交互项进行层次回归分析, 结果见表 5。回归分析显示, 人机信任与系统透明度的交互项显著正向影响 人 机 协 同 决 策 质 量 ($M8, \beta = 0.069, p < 0.05$), 验证了假设 H_5 。

表 5 系统透明度调节效应检验

Table 5 System transparency moderating effect test				
变量类型	变量	人机协同决策质量		
		模型 6	模型 7	模型 8
控制变量	性别	-0.113 *	-0.050	-0.046
	年龄	0.003	0.025	0.021
	工龄	-0.163 *	-0.037	-0.035
	学历	-0.060	-0.032	-0.028
	行业	-0.046	-0.028	-0.024
中介变量	人机信任	—	0.817 ***	0.815 ***
调节变量	系统透明度	—	0.027	0.026
交互项	人机信任×系统透明度	—	—	0.069 *
—	R^2	0.042	0.683	0.687
—	ΔR^2	0.042	0.641	0.005
—	F	3.806 **	132.775 ***	118.432 ***
—	ΔF	3.806 **	436.117 ***	6.403 ***

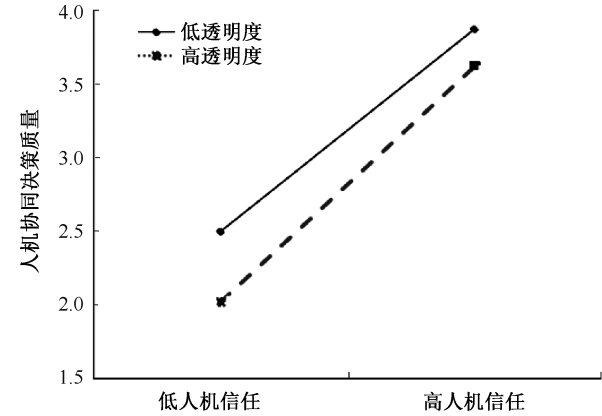


图 2 系统透明度的调节效应

Fig. 2 Moderating effect of system transparency

为更直观呈现系统透明度的调节效应,绘制简单斜率图(图 2)。结果显示:系统透明度较高时,人机信任对人机协同决策质量的正向影响更强,表明系统透明度对该关系具有正向调节作用, H_5 得到验证。

为进一步检验系统透明度的调节效应,采用简单斜率显著性检验方法分析系统透明度在不同水平下人机信任对人机协同决策质量的影响强度及其显著性。结合系统透明度高低水平的效应值(表 6),当系统透明度处于低水平时,人机信任对人机协同决策质量的效应估计值为 0.817 2 (95% CI = [0.741 9, 0.892 5]);当系统透明度处于高水平时,效应估计值为 0.938 9 (95% CI = [0.870 1, 1.007 6])。两者置信区间均不包含 0,表明在高和 low 系统透明度的时候,人机信任和人机协同决策质量的关系均显著。假设 H_5 得到支持。

表 6 调节效应(Bootstrap 法) 检验结果

Table 6 Test results for moderating effects(Bootstrap)				
系统透明度	效应值	标准误	95% 置信区间	
			低限	高限
低系统透明度	0.817 2	0.383	0.741 9	0.892 5
高系统透明度	0.938 9	0.035 0	0.870 1	1.007 6

4 AI 焦虑与人机信任、系统透明度研究讨论

4.1 AI 焦虑与人机协同决策质量

证实 AI 焦虑显著降低人机协同决策质量,企业若不及时识别与干预,可能削弱安全控制与风险防范。该效应具有情境差异:高风险行业因 AI 深度嵌入预警、联锁与应急环节,系统不确定性更易诱发焦虑,导致不信任预警、绕开联锁或延迟采纳建议,显著削弱协同安全效能;低风险行业 AI 多用于流程优化,焦虑主要体现为适应与效率担忧,对安全影响相对有限且可通过培训缓解。因此,高风险行业应重点提升系统透明度、强化培训与安全文化建设,低风险行业则侧重增强员工对 AI 的理解与接受度;组织层面可建立 AI 心理监测与干预机制并开展针对性培训,以保障协同决策稳定有效。

4.2 人机信任

人机信任在 AI 焦虑与人机协同决策质量之间发挥关键中介作用,拓展了 UMT 在安全管理领域的解释边界。较高的人机信任可缓冲 AI 焦虑带来的认知干扰,提升协同决策的准确性与一致性,从而降低安全判断失误风险。鉴于信任既受 AI 系统性能与透明度影响,也受组织沟通、管理支持与安全文化塑造,企业推进 AI 落地应同步构建可信技术环境与协作机制,以保障关键安全决策稳定可靠。

4.3 系统透明度

系统透明度显著调节人机信任对人机协同决策质量的作用,拓展了 UMT 关于情境变量调控的应用边界,也凸显其在高风险决策环境中的治理价值。实践中,应强化 AI 透明性建设,公开关键决策流程、解释推理逻辑并完善风险提示,以减少误解与事故风险,提升 AI 在安全管理中的可靠应用。

5 结 论

1) 从安全管理视角出发,以 UMT 为理论基础,本文厘清了 AI 焦虑在安全管理实践中的作用机制与演化路径,即 AI 焦虑通过降低人机信任,进而显

著削弱人机协同决策质量,并对企业安全风险管控与安全决策有效性产生负面影响。

2) 系统透明度在人机信任与人机协同决策质量之间发挥调节作用,当透明度较高时,员工更易理解 AI 决策逻辑、运行机制及潜在风险,从而强化人机信任对协同决策质量的促进效应;当透明度较低时,该效应相对减弱。

3) 研究虽在一定程度上控制共同方法偏差,但仍有局限,一是多行业样本提升代表性,但行业在技术依赖、风险环境与组织文化上的差异可能影响结论的稳定性与普适性;二是尽管采用多阶段问卷,数据本质仍偏横截面,因果推断的稳健性受限。未来

可聚焦特定行业,并结合行为实验、眼动追踪或脑电等方法开展实时测量,以刻画 AI 焦虑对协同决策的动态影响与心理机制。

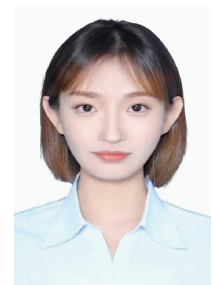
4) 鉴于研究设计阶段未系统区分行业风险水平、AI 技术类型及关键技术特征,未来可从 3 方面拓展:①开展行业异质性分析,按风险等级与技术依赖度分组比较以提升外部效度;②细分 AI 技术类型,检验其对焦虑与决策行为的差异化影响;③结合可解释性等特征扩展系统透明度维度,并引入行为指标进一步揭示其调节机制,从而增强结论的情境适用性与实践指导价值。

参 考 文 献

- [1] 王大龙,王冰山,曹睿,等. 煤矿安全管理行为交互机制与模型[J]. 中国安全科学学报, 2025, 35(5): 32-38.
WANG Dalong, WANG Bingshan, CAO Rui, et al. Interaction mechanism and model of coal mine safety management behaviors[J]. China Safety Science Journal, 2025, 35(5): 32-38.
- [2] 王秉,周佳胜,史志勇. 数智赋能安全态势感知与塑造模型[J]. 图书馆论坛, 2024, 44(8): 1-7.
WANG Bing, ZHOU Jiasheng, SHI Zhiyong. A model for safety & security situational awareness and shaping empowered by digital intelligence [J]. Library Tribune, 2024, 44(8): 1-7.
- [3] 叶进杰,李景升,史志勇. 数智赋能情报驱动的关键信息基础设施安全韧性建设模型研究[J]. 情报杂志, 2025, 44(7): 61-68, 76.
YE Jinjie, LI Jingsheng, SHI Zhiyong. Research on the model of safety resilience construction of critical information infrastructure driven by intelligence empowered by digital intelligence[J]. Journal of Intelligence, 2025, 44(7): 61-68, 76.
- [4] JOHNSON D G, VERDICCHIO M. AI anxiety[J]. Journal of the Association for Information Science and Technology, 2017, 68: 2 267-2 270.
- [5] SUSENO Y, CHANG C, HUDÍK M, et al. Beliefs, anxiety and change readiness for artificial intelligence adoption among human resource managers: the moderating role of high-performance work systems [J]. The International Journal of Human Resource Management, 2021, 33: 1 209-1 236.
- [6] LI Jian, HUANG Jinsong. Dimensions of artificial intelligence anxiety based on the integrated fear acquisition theory[J]. Technology in Society, 2020, 63: DOI:10.1016/J.TECHSOC. 2020. 101410.
- [7] 郭娟,朱晓妹,李姿颖,等. 人工智能技术应用对员工心理状态的影响分析[J]. 中国人事科学, 2021(11): 50-58.
GUO Juan, ZHU Xiaomei, LI Ziyang, et al. Analysis on influence of application of artificial intelligence technology on psychological state of employees[J]. Chinese Personnel Science, 2021(11): 50-58.
- [8] 刘小禹,余彩婷. 悲欣交集:数字技术与员工情绪[J]. 外国经济与管理, 2024, 46(6): 134-152.
LIU Xiaoyu, YU Caiting. A mixture of joy and sorrow: digital technologies and employee emotions[J]. Foreign Economics & Management, 2024, 46(6): 134-152.
- [9] 杜佩佳,刘生敏. 人工智能应用对员工心理的双刃剑影响[J]. 运筹与模糊学, 2024, 14(4): 541-547.
DU Peijia, LIU Shengmin. The double-edged sword effect of artificial intelligence application on employees' psychology[J]. Operations Research and Fuzzy Systems, 2024, 14(4): 541-547.
- [10] 蒋建武,龙吟寰,胡洁宇. 工作场所人工智能应用对员工影响的元分析[J]. 心理科学进展, 2024, 32(10): 1 621-1 639.
JIANG Jianwu, LONG Hanhuan, HU Jieyu. A meta-analysis of the effects of artificial intelligence application in the workplace on employees[J]. Advances in Psychological Science, 2024, 32(10): 1621-1639.
- [11] CHEN Min, NIKOLAIDIS S, SOH H, et al. Trust-aware decision making for human-robot collaboration [J]. ACM Transactions on Human-Robot Interaction (THRI), 2018, 9: 1-23.
- [12] 胡保亮,梅婕,张素平. 人机协同决策:文献综述[J]. 信息与管理研究, 2024, 9(增1): 51-64.
HU Baoliang, MEI Jie, ZHANG Suping. Human-AI collaborative decision-making: a literature review[J]. Journal of Information and Management, 2024, 9(Z1): 51-64.
- [13] 田晓明,段锦云,孔瑜. 不确定管理理论:核心概念、内容介绍及未来展望[J]. 心理研究, 2012, 5(3): 70-77.

- TIAN Xiaoming, DUAN Jinyun, KONG Yu. A review of uncertainty management theory: conception, introduction and future directions[J]. *Psychological Research*, 2012, 5(3): 70-77.
- [14] BRASHERS D. Communication and uncertainty management [J]. *Journal of Communication*, 2001, 51: 477-497.
- [15] GARY Z. The influence mechanism of AI technology learning anxiety in the human-computer collaboration context: the moderating effect of uncertainty avoidance and the mediating effect of self-efficacy [J]. *The EUrASEANs: Journal on Global Socio-Economic Dynamics*, 2023, 6(43): 170-180.
- [16] SETHUMADHAVAN A. Trust in artificial intelligence [J]. *Ergonomics in Design: the Quarterly of Human Factors Applications*, 2018, 27; DOI:10.1177/1064804618818592.
- [17] 罗方禄, 罗吴赞. 新质生产力安全发展的内涵及能力建设[J]. *中国安全科学学报*, 2024, 34(10): 8-16.
- LUO Fanglu, LUO Wuzan. The connotation and capacity building of the safe development of new quality productive forces[J]. *China Safety Science Journal*, 2024, 34(10): 8-16.
- [18] MAYER R C, DAVIS J H, SCHOORMAN F D. An integrative model of organizational trust [J]. *Academy of Management Review*, 1995, 20(3): 709-734.
- [19] MYERS C D, TINGLEY D. The influence of emotion on trust [J]. *Political Analysis*, 2016, 24: 492-500.
- [20] MATSUNAGA M. Uncertainty management, transformational leadership, and job performance in an AI-powered organizational context [J]. *Communication Monographs*, 2022, 89(1): 118-139.
- [21] 何贵兵, 陈诚, 何泽桐, 等. 智能组织中的人机协同决策: 基于人机内部兼容性的研究探索[J]. *心理科学进展*, 2022, 30(12): 2 619-2 627.
- HE Guibing, CHEN Cheng, HE Zetong, et al. Human-AI collaborative decision-making in intelligent organizations: an exploratory study based on internal human-AI compatibility[J]. *Advances in Psychological Science*, 2022, 30(12): 2619-2627.
- [22] TOREINI E, AITKEN M, COOPAMOOTOO K, et al. The relationship between trust in AI and trustworthy machine learning technologies [C]. *Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency (FAT'20)*, 2020: 272-283.
- [23] KALTENBACH E, DOLGOV I. On the dual nature of transparency and reliability: rethinking factors that shape trust in automation [J]. *Proceedings of the Human Factors and Ergonomics Society Annual Meeting*, 2017, 61: 308-312.
- [24] 胡艺龄, 陈彦君. 人机协作下的共享心智模型研究[J]. *现代教育技术*, 2024, 34(1): 64-72.
- CHEN Yanjun, HU Yiling. Research on shared mental model under human-machine collaboration [J]. *Modern Educational Technology*, 2024, 34(1): 64-72.
- [25] BHASKARA A, SKINNER M, LOFT S. Agent transparency: a review of current theory and evidence [J]. *IEEE Transactions on Human-Machine Systems*, 2020, 50(3): 215-224.
- [26] GULATI S, SOUSA S, LAMAS D. Design, development and evaluation of a human-computer trust scale [J]. *Behaviour & Information Technology*, 2019, 38(10): 1 004-1 015.
- [27] SHAMIM S, ZENG J, SHARIQ S M, KHAN Z. Role of big data management in enhancing big data decision-making capability and quality among Chinese firms: a dynamic capabilities view[J]. *Information & Management*, 2019, 56(6): DOI: 10.1016/j.im.2018.12.003.
- [28] POPE-DAVIS D, TWING J S. The effects of age, gender, and experience on measures of attitude regarding computers[J]. *Computers in Human Behavior*, 1991, 7: 333-339.
- [29] MARGRETT J A, MARSKE M. Gender differences in older adults' everyday cognitive collaboration [J]. *International Journal of Behavioral Development*, 2002, 26(1): 45-59.
- [30] INKPEN K, CHAPPIDI S, MALLARI K, et al. Advancing human-AI complementarity: the impact of user expertise and algorithmic tuning on joint decision making [J]. *ACM Transactions on Computer-Human Interaction*, 2022, 30: 1-29.
- [31] ZHANG Fenghua, XIAO Leifeng, GU Ruolei. Does gender matter in the relationship between anxiety and decision-making? [J]. *Frontiers in Psychology*, 2017, 8: DOI:10.3389/fpsyq.2017.02231.

作者简介: 牛莉霞 (1983—), 女, 山西吕梁人, 博士, 教授, 主要从事安全管理与人因工程方面的研究。
E-mail:327334231@qq.com。



李波 (1998—), 女, 辽宁阜新人, 硕士研究生, 研究方向为安全管理与组织行为学。E-mail:1427580275@qq.com。