

人工智能生成内容(AIGC)信息回避影响 机制研究——AI 身份威胁的调节作用

张 玥* 姜冠岐 李璐含

(南京大学信息管理学院, 江苏 南京 210023)

摘 要: [目的/意义] 本文研究了用户对人工智能生成内容(AIGC)信息回避的成因与作用机制, 为用户信息行为引导、AIGC 技术优化和 AIGC 产品迭代提供理论参考。[方法/过程] 基于认知—情感—意愿(Cognition-Affect-Conation, C-A-C)框架, 构建信息可控度、信息透明度和 AI 焦虑影响信息回避的理论模型, 通过虚拟实验方式获取样本数据, 并利用偏最小二乘法结构方程模型(PLS-SEM)对数据及模型进行分析与验证。[结果/结论] 结果表明, 用户对人工智能生成内容(AIGC)的信息回避行为在不同 AIGC 信息透明度与可控度之间存在显著差异, AI 焦虑与信息回避行为有正向关系。研究发现, AI 身份威胁对 AIGC 信息不确定性与 AI 焦虑之间的关系有正向调节作用。

关键词: 人工智能(AI); 人工智能生成内容(AIGC); 信息回避; AI 身份威胁; AI 焦虑

DOI:10.3969/j.issn.1008-0821.2025.04.006

[中图分类号] G252.0 [文献标识码] A [文章编号] 1008-0821 (2025) 04-0060-14

Research on the Influencing Mechanism of Users' Avoidance Intention towards AIGC Information——The Moderating Role of AI Identity Threat

Zhang Yue* Jiang Guanqi Li Luhan

(Department of Information Management, Nanjing University, Nanjing 210023, China)

Abstract: [Purpose/Significance] This paper focuses on investigating the causes and mechanisms behind users' intention to avoid artificial intelligence-generated content(AIGC) information, aiming to provide theoretical guidance for user information behavior, optimization of AIGC technology, and iteration of AIGC products. [Method/Process] Drawing upon the C-A-C framework, this study constructed a theory model to examine the influence of controllable degree of information, transparency of information, and AI anxiety on information avoidance. The study collected sample data through virtual experimental design, and analyzed data using partial least square structural equation modeling(PLS-SEM). [Result/Conclusion] The findings reveal significant differences in users' avoidance intention towards AIGC information based on its transparency and controllability. Additionally, AI anxiety exhibits a positive relationship with information avoidance behavior. Notably, it is important to acknowledge that AI identity threat positively moderates the relationship between AIGC information uncertainty and AI anxiety.

Key words: artificial intelligence(AI); artificial intelligence generated content(AIGC); information avoidance; AI identity threat; AI anxiety

收稿日期: 2024-04-18

基金项目: 国家社会科学基金项目“信息茧房催化下群体性迷失形成机制与治理策略研究”(项目编号: 23BTQ090); 教育部人文社会科学研究基金项目“移动健康管理中信息回避形成机制与演化规律研究”(项目编号: 21YJC870020); 江苏省社会科学基金项目“面向移动健康管理的信息回避行为模式识别与演化规律研究”(项目编号: 21TQB003)。

作者简介: 姜冠岐(1998-), 女, 硕士研究生, 研究方向: 用户信息行为。李璐含(2002-), 女, 硕士研究生, 研究方向: 用户信息行为。

通信作者: 张玥(1982-), 女, 教授, 博士生导师, 研究方向: 用户信息行为。

人工智能生成内容(Artificial Intelligence Generated Content, AIGC)是以人工智能(Artificial Intelligence, AI)技术为基础,基于各种形式的人工指令和指导生成的具有相应特征的内容^[1]。《2024 年 AIGC 应用层十大趋势》白皮书显示,2023 年全球企业在 AIGC 解决方案上投资超过 160 亿美元,预计到 2027 年这一数字将膨胀约 10 倍,年复合增长率超过 70%^[2]。各类 AIGC 应用和工具不断涌现,如以文心一言、IBM Watson 为代表的文本写作工具;以 Midjourney、Amazon Rekognition 为代表的图像生成工具;以 WaveNet、Synthesia 为代表的音视频生成工具等。随着 AIGC 大语言模型的不断优化迭代,国内外也研发出大量的 AIGC 工具,为金融、科技、教育、医疗等多个行业的发展提供了技术支撑,进而为企业节约成本,给用户提供了便利。

与此同时,各类隐私泄露、版权争议、算法偏见等 AIGC 应用乱象及负面问题亦层出不穷。这些风险不仅触动了公众对 AIGC 技术安全的敏感神经,更在个人用户中触发了对 AIGC 的抵触情绪,甚至催生了各种回避行为。例如,2023 年国内社交平台 LOFTER 上线了一款允许用户通过关键词生成 AI 头像的“头像生成器”,迅速引起很多用户的不满和抵制。同年,外媒 NewsGuard 发布的数据报告显示,已追踪到约 739 个文字和图片内容呈现鲜明“GPT 风”的 AIGC 新闻网站,这些 AIGC 内容大量涌入了 Twitter、Youtube、TikTok 等主流社交媒体,导致人们对 AIGC 内容产生抵制和信息回避行为^[3]。

近年来,信息回避受到信息行为领域高度关注,图书情报、计算机科学、管理科学、传播学等学科的学者针对用户生成内容(User Generated Content, UGC)模式下的信息回避影响因素展开了大量实证研究。理解用户对 AIGC 信息回避行为的影响机制对深入探索 AIGC 技术的复杂影响、促进其在用户信息行为领域发挥积极作用,具有重要研究意义。这不仅有助于揭示用户在面对 AIGC 生成内容时的心理动态和行为倾向,而且对优化 AIGC 内容的设计、提升用户接受度、构建和谐的信息生态环境也有极大帮助。通过分析用户回避 AIGC 信息的影响因素与内在作用关系,可以更准确地评估 AIGC 在用户日常生活中的融入程度、潜在风险、挑战,以及长

期的信息交互趋势。此外,研究 AIGC 信息回避影响机制能够为技术开发者、内容提供者和政策制定者提供参考,指导他们在保障用户权益、维护信息真实性和促进技术健康发展之间找到平衡。身份威胁概念在不同领域的研究中被进一步解释并衍生出了新的概念外延,本研究聚焦 AIGC 视角下的身份威胁,即用户根据以往的 AIGC 服务或技术使用体验,对当前身份自我信念的威胁预判。AI 身份威胁(AI Identity Threat)是指人工智能对身份个人自我信念伤害的预期,这是由 AI 的使用引起的,而它所适用的实体是所有使用 AI 技术的个人用户。因此,深入探索用户 AIGC 信息回避的影响因素及 AI 身份威胁在其中的内在作用机制,将有助于更好地理解用户的需求和行为,从而完善 AIGC 的设计和功能,提升用户体验,促进 AIGC 技术的持续发展和广泛应用。

1 研究综述

1.1 图书情报领域 AIGC 相关研究综述

目前,图书情报领域关于 AIGC 的研究主要集中在数智融合场景中 AIGC 技术的赋能应用和 AIGC 用户研究两个维度:对于 AIGC 技术的应用研究重点关注了 AIGC 技术在智慧图书馆^[4-5]、情报智能服务^[6]、信息资源管理工具^[7]、企业数智管理系统^[6]等领域的 AIGC 技术应用前景,研究 AIGC 技术特征、要素及发展阶段,基于互联网演进,探讨其网络形态、内容生产、人机交互及资源组织基础,分析数据、模型、空间赋能的技术特征,并探讨 AIGC 对信息资源管理的实质影响^[6,8-9];从用户研究维度出发,当前研究集中在探究个体对 AIGC 的认知、态度以及情感等主观评价的影响因素及其作用机制。这意味着当前的研究不仅关注 AIGC 本身的技术特性,更着眼于这些特性如何影响用户的认知和情感。如对 AIGC 可信度的评估:通过范性分析构建了以信息源可信度、交互方式可信度、社会行动者可信度以及算法隐喻可信度 4 个维度的 AIGC 可信度评价研究框架^[10];对人工智能生成内容的用户采纳意愿研究^[11-12],以及更进一步研究用户对不同类型 AIGC 的持续使用意愿影响因素等^[13]。

综上所述,当前国内外关于 AIGC 的研究主要涉及 AIGC 技术特征、不同 AIGC 技术模型的实践

应用和用户在 AIGC 使用过程中的接纳意愿及采纳行为。虽然 AIGC 在应用中展现了许多正面的积极效应,但部分学者也开始意识到新技术背后的消极情绪和行为,包括 AI 焦虑^[14]、算法歧视^[15]、算法厌恶^[16-17]、中辍行为^[18-19]等,现实中对 AIGC 内容产生回避态度和行为的情况更是广为存在。因此,需要深入探究用户回避行为背后的原因及影响,这不仅能更加理解用户的行为模式,还将为 AIGC 赋能产品、AIGC 大语言模型以及多模态技术的发展提供宝贵的启示和指导。

1.2 信息回避相关研究综述

信息回避(Information Avoidance)这一现象的提出源于 20 世纪心理学中关于选择性注意的研究^[20],指人们避免或延迟获取可用却不想要信息的特殊信息行为^[21]。学界按照回避的程度,将信息回避划分为完全性信息回避和选择性信息回避^[22-23]。

目前信息回避的学术研究主要涉及信息因素、个体因素和环境因素 3 个层面:信息因素关注信息本身的特点,有研究证实了信息过载、信息媒介依赖和信息相似性等信息因素的影响^[24-25],也有相关研究专门探究信息质量、信息成本以及信息隐私等信息因素对运动损伤人群健康信息回避模式的影响^[26];个体因素是指个人的心理、认知和情感状态对信息回避行为的影响^[27],其中认知和情感因素发挥着重要作用,有研究证明社会比较对用户认知失调和信息回避均有重要影响^[28];情境因素则涉及个体所处的社会、文化背景和特殊情境对信息回避行为的影响,如专对新冠感染情境下的信息规避行为因素及其相互关系^[27]。此外,多位国内外学者在信息回避的相关影响因素研究中运用认知—情感—意图框架,从该角度构建信息回避模型,将信息回避意图作为因变量,探究环境因素、个体因素、信息因素对信息回避意图的影响路径^[24,29]。

综上所述,目前信息回避相关研究主要包括信息因素、个人因素和情境因素,涉及信息过载^[30]、信息质量^[31]、信息相关性^[32]、认知冲突^[33]、感知控制^[34]等相关主题,同时该领域学者多从不同场景切入,讨论学术信息回避^[24]、社交信息回避^[22]和健康信息回避^[26,28]的影响因素。但是,目前大多数研究仍聚焦于 UGC 模式下的信息回避行为,而随

着技术不断革新,AIGC 应用和创新模式的不断涌现与发展,其在内容创作门槛、创作模式和创作效率等方面所带来的深刻变革,已逐渐改变了用户的信息行为,如 AIGC 技术在信息生成、传播和利用方面的独特性,用户可能会对其产生不同的认知评价,从而影响他们的信息回避行为。因此,迫切需要从 AIGC 的实际生成、寻求、利用等场景中,深入探索用户的认知评价与用户信息回避行为之间的联系,以期为用户与 AIGC 的互动以及 AIGC 的进一步发展提供启示。

2 研究模型与理论假设

2.1 研究模型

2.1.1 理论基础

1) 认知—情感—意图框架

认知—情感—意图(或意动)框架(Cognitive-Affective-Conative, C-A-C)由认知心理学者 Fishbein M 等^[35]首次提出,该理论框架认为个体对事件的认知评价塑造了他们的情感反应,随后影响了他们的行为。Lazarus S R^[36]在该理论的基础上对每个阶段进行了解释补充,提出 C-A-C 框架包括 3 个阶段:认知过程、情感反应和行为意图。C-A-C 框架植根认知评价理论和情感认知动机关系理论,主张个体对事件的认知评价塑造了个体的情感反应,进而影响其意动^[37]。

当前 C-A-C 框架多应用于解释电子商务情境下消费者决策行为的影响机制。例如,Li J 等^[38]提出了包括客户对 Airbnb 体验的评价(认知)、情感反应(情感)和忠诚行为(意动)的理论模型,用以了解用户决策行为的影响因素。该框架也被证实理解消费者对 IT 服务创新的接受方面是有用的^[39],有学者基于该框架挖掘服务机器人的属性(即认知评价),以反映服务机器人的技术和社会方面对用户的情感反应中的影响^[40]。此外,C-A-C 框架也在信息回避的相关研究中得到运用^[24,29]。

2) 不确定性管理理论

不确定性管理理论(Communication and Uncertainty Management Theory)是由美国学者 Brashers E D^[41]在 2000 年提出,最早用于研究健康风险信息沟通的不确定性管理。该理论指出当人们认为信息不确定时,倾向于采取抵抗或者回避行为,以避免进入焦虑和纠结的情绪之中。亦有研究表明,不

确定性的认知是负面情绪重要来源^[42]，目前很多学者基于该理论探究了健康信息^[28,32,43]、学术信息^[24]、社交信息^[22,33]等多种情景下的不确定性认知因素对于信息回避行为的直接效应。

3) 情绪认知理论

情绪认知理论 (Cognitive Theory of Emotion) 认为情绪是人类在面对外部刺激时所感受到的一种主观体验，情绪的产生是通过对外部刺激的认知评估来实现的^[44]。因此，在相同的情境下，不同个体可能呈现出不同的情绪反应^[45]。研究表明，个人过往经验，特别是困难性、威胁性事件，将会成为其预判和评估当前的情绪体验的重要依据^[46]。

2.1.2 概念模型构建

本文基于 C-A-C 理论模型，构建从认知 (Cognitive) 到情绪 (Affective) 再到意图 (Conative) 的过程。基于不确定性管理理论，本文将不确定性认知作为用户在 AIGC 信息回避过程中的认知因素。由于 AIGC 技术尚在摸索阶段且相关法律体系尚未确立，导致使用 AI 生成的海量信息出现不透明、不可解释、不公平、泄露隐私等问题，致使用户在使用的过程中对 AIGC 产生很多不确定性的认知。相关研究指出，信息透明度^[47]和感知控制^[48]是用户信息不确定性的重要组成部分，因此将这些因素引入 C-A-C 框架，作为不确定性认知变量进行探讨。同时，不确定性认知会影响个体内部的情绪及情感反应从而引发焦虑^[21]。先前的研究已经表明，积极的情绪会促使个体接近信息，而消极的情绪 (如焦虑) 则会引发个体的回避行为^[25]。因此，本文构建不确定性认知—焦虑情绪—信息回避行为的理论框架。

情绪认知理论指出，威胁感在认知到情绪过程中起到调节作用^[49-50]，正如今天很多人感到“机器越来越像人，而人却远离其本身，反而更像一架机器了”^[51-52]，这样的感受凸显用户在对 AIGC 产品使用过程中的身份威胁感，因此，本文也将 AI 身份威胁纳入研究框架，作为认知情绪的调节部分。

因此，在 C-A-C 框架和不确定性管理理论、情绪认知理论相结合的基础上，本文概念模型如图 1 所示。

2.2 理论假设

在 C-A-C 的整体框架和不确定性管理理论、情

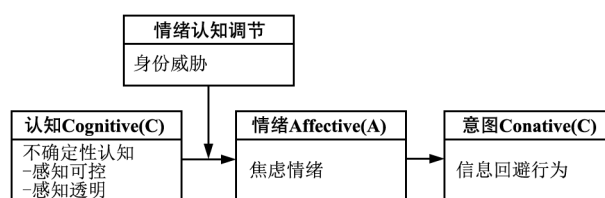


图 1 基于 C-A-C 框架的概念模型

Fig. 1 Theory Model Based on C-A-C Framework

绪认知理论的基础上，本文依次提出以下假设。

2.2.1 不确定性认知与 AI 焦虑 (C→A) 关系假设

在信息行为学领域已有大量研究表明，透明度和控制度是解释不确定性认知的重要因素^[53-56]，因此将从这两个方面提出相关假设。

信息透明度被定义为信息可以被人们访问和查看的程度^[57]。感知透明是不确定性认知的组成内涵之一，信息透明度的降低会引发不确定性认知的减弱^[58]。在人工智能领域，AIGC 透明度反映了技术 (人工智能) 的潜在机制和内在逻辑对用户的明显程度，其在信任建立过程中发挥着重要作用^[59]。如果尽可能让人们了解 AIGC 形成的基本原理、AIGC 运用的算法以及其来源，会提高用户对 AIGC 的感知透明度，从而降低信息不确定性所引发的焦虑情绪。因此，本研究提出以下假设：

H1：个体认知的 AIGC 透明度对 AI 焦虑情绪有负向影响

感知控制理论将个体对自身行为所施加影响能力的信念界定为感知控制^[60-61]。在用户使用 AIGC 技术的过程中，感知控制则是指个体对技术应用过程中相关信息的感知可控度^[62]。过往研究已证明，通过完善平台架构提高互联网终端用户对技术的感知可控度，可以有效缓解用户的 AI 焦虑^[63]，较高的 AIGC 可控度水平会促使用户产生更高水平的信任^[64]。相反，Schweitzer F 等^[65]的研究表明，一旦人们感受到在与智能设备互动时无法自我控制，焦虑和担忧就会产生。因此，本研究提出以下假设：

H2：AIGC 可控度对 AI 焦虑情绪有负向影响

2.2.2 不确定性认知与信息回避 (C→C) 关系假设

根据上文所述，当 AIGC 形成的潜在机制和内在逻辑能够使用户理解，会提高其对 AIGC 的感知透明度，降低信息不确定性所引发的焦虑情绪，最终有效减少信息回避。因此，本研究提出以下假设：

H3：个体认知的 AIGC 透明度对 AIGC 信息回

避有负向影响

已有研究表明,人智交互过程中,掌控感可能影响人们对日常生活和工作生活的自我效能,进而引发信息回避行为, Park C S 等^[66]专门将 AIGC 相关技术服务的规避行为描述为提高控制感的策略。因此,本研究提出以下假设:

H4: AIGC 可控度对 AIGC 信息回避有负向影响

2.2.3 AI 焦虑与信息回避(A→C)关系假设

心理学研究表明,焦虑情绪指的是人们对环境不确定性感到紧张和烦恼的心理状态。用户焦虑情绪是信息过载带来的一种新社会现象,在信息行为领域中备受关注。很多研究从不同视角,如健康焦虑^[28]、技术焦虑^[40]、学术焦虑^[24]证明了焦虑对信息回避的直接影响。因此,本文将用户接触 AI 所产生的各类信息内容的过程中感到的烦躁和不安情绪界定为 AI 焦虑,并提出以下假设:

H5: AI 焦虑情绪对 AIGC 信息回避有正向影响

2.2.4 AI 身份威胁对认知和情绪之间的调节作用假设

身份威胁(Identity Threat)是指对个人身份相关的自我信念所产生的潜在危害的预期^[67]。这样的一种威胁感会让用户处在身份矛盾和自尊丧失时的焦虑和痛苦中^[68]。基于此, Mirbabaie M 等^[69]提出, AI 身份威胁(AI Identity Threat),并界定其是由 AI 使用所引起的,对个人身份相关的自我信念伤害的预期。根据情绪认知理论和身份威胁概念,当输入意义与认同标准的不一致(即认同不验证)出乎意料或难以纠正时,就会产生更强烈的负面情绪。通过将身份威胁整合到情绪认知理论中,本文认为与身份威胁高的个体相比,身份威胁低的个体对身份自我信念的风险预判更加不敏感,因此,他们对于处理 AIGC 信息的不确定性认知持有更乐观的态度。这些人坚信他们在此过程中的表现良好,因此对处理 AIGC 不确定性认知的评估不足,从而更不容易产生 AI 焦虑。此外,由于对身份的自我信念,高身份威胁的个体不太可能通过降低不确定性管理标准来与情境中的消极意义保持一致从而对不确定性做出让步。如果不调整不确定性管理的标准,与持有低 AI 身份威胁的个体相比,高 AI 身份威胁的个体更难以克服认知—情感—意图框架中的“认知—情绪”回路中的不确定性,呈现出更强的 AI 焦虑

程度。因此,结合情绪认知理论,本文认为与 AI 身份威胁高的个体相比, AI 身份威胁低的个体对身份自我信念的风险预判更加不敏感,他们对于处理 AIGC 信息的不确定性认知持有更乐观的态度,从而更不容易产生 AI 焦虑。因此,本研究提出以下假设:

H6: AI 身份威胁会调节 AIGC 透明度对 AI 焦虑情绪的影响,当 AI 身份威胁增强时,这种影响会增强

H7: AI 身份威胁会调节 AIGC 可控度对 AI 焦虑情绪的影响,当 AI 身份威胁增强时,这种影响会增强

2.2.5 假设模型构建

基于图 1 的概念模型,本研究通过对前人研究结论进行梳理总结,提出了关于 AIGC 信息回避影响机制的 7 个假设,基于认知—情感—意图框架,分别引入了 AIGC 透明度、AIGC 可控度、AI 焦虑、AI 身份威胁,最终构建如图 2 所示的研究假设模型。

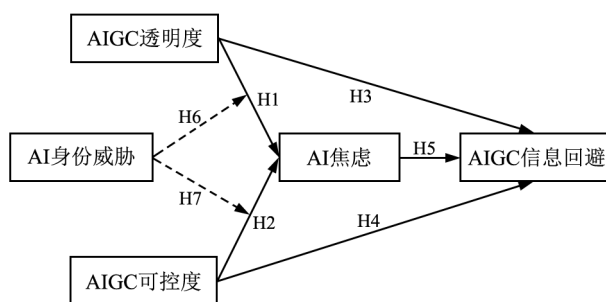


图 2 研究假设模型

Fig. 2 Research Hypothesis Model

3 数据收集

Precedence Research^[70]发布的《AIGC 产业报告》中显示,2023 年全球日均活跃用户 TOP50 的 AI 应用中,文本生成类 AIGC 应用的数量占据 50% 左右,流量占据 80% 以上,因此,本研究通过 Axure RP 8.1 原型设计软件模拟绝大部分用户都接触过的信息搜索情境中的文本生成类 AIGC 任务,并结合在线调研进行最终的数据收集。搜索问题选自微博、抖音、今日头条三大社交平台的热门日常生活类话题(如“小白鞋脏了怎么清洗?”等)。搜索结果中的 AIGC 信息使用 ChatGPT、文心一言、New-Bing 等各类 AIGC 平台对任务问题真实回答并进行人工矫正,剔除了与问题不相关的内容以及难以理解的内容。

为了让用户更清楚地理解 AI 搜索技术原理、AIGC 数据来源及使用注意事项,依据前述 AIGC 透明度概念设置了“参考内容”“用户须知”选项。“参考内容”包括参考内容的标题、作者及来源网址,“用户须知”告知被试者平台运用大语言模型技术结合大数据归纳总结形成最终结果。在可控度设计中,根据前述 AIGC 可控度的概念设置了“语言风格”“语言种类”选项。“语言种类”包括中文和英文两种,英文信息通过谷歌翻译将中文材料进行直译得到,并经过人工矫正,确保中英文信息表述一致。“语言风格”包括书面化风格和口语化风格两种,两种信息均来源于 ChatGPT 的回答。为排除 AIGC 信息内容对实验结果的干扰,本研究确认了信息量的一致性以及两种风格的差异化,AIGC 信息均经过相关领域专家的检验。两种可控度触达控件在 AIGC 搜索结果信息下方,只需点击“语言种类”及“语言风格”下拉列表即可控制这些选项,允许用户根据自己需求和偏好进行调整,并操作化地定义和区分组间可控度水平,图 3 展示了虚拟平台设计框架。在正式获取数据前,招募 32 名参与者,随机分配到透明度和可控度的高低水平组,进行操纵性检验。结果如表 1 和表 2 所示,参与者在虚拟实验情境下能够明显感受到 AIGC 信息透明度和可控度的差异,操纵性检验成功。



图 3 虚拟平台 AIGC 任务页面示意图
Fig. 3 AIGC Task Page Diagram on Virtual Platform

表 1 AIGC 透明度预实验 ANOVA 分析结果

Tab. 1 ANOVA Analysis Results of AIGC Transparency

Pre-experiment					
	平方和	df	均方	F 值	显著性
组间	112.431	1	112.431	153.716	0
组内	29.86	31	0.702		
总计	142.291	32			

表 2 AIGC 可控度预实验 ANOVA 分析结果

Tab. 2 ANOVA Analysis Results of AIGC Controllability

Pre-experiment					
	平方和	df	均方	F 值	显著性
组间	82.474	1	82.474	104.395	0
组内	28.618	31	0.806		
总计	111.092	32			

数据收集通过线上目的性抽样及线下滚雪球抽样相结合的方法,共招募 56 名参与者,基于本研究选题,排除没有 AIGC 使用经验的参与者,最终参与虚拟实验和调研问卷的共计 40 人。研究样本描述性统计如表 3 所示,性别分布上,男性受试者占比 55%,略高于女性受试者的 45%。年龄结构中,18~25 岁的年轻群体占据了绝大多数,占比高达 77.5%,这与 Questmobile^[71]发布的《2024 生成式 AI 及 AIGC 应用洞察报告》中提到的 AIGC 用户群体特征相吻合。教育背景方面,本研究中研究生学历的受试者占比达 52.5%,显示出较高的教育水平。职业领域中,IT/互联网行业的从业者占比最高,这与报告中 AI 产业数据中 AIGC 用户偏好浏览教育科技相关内容的趋势相一致,进一步验证了本研究样本的代表性。综上所述,本研究的样本选择与国内 AIGC 用户的实际特征相符,为研究结果

的准确性和可靠性提供了坚实的基础。

表 3 样本特征描述性统计分析
Tab. 3 Descriptive Statistical Analysis of Sample Characteristics

样本特征	选 项	频率	百分比 (%)
性 别	男	22	55.00
	女	18	45.00
年 龄	18 岁以下	0	0.00
	18~25 岁	31	77.50
	26~30 岁	5	12.50
	31~40 岁	4	10.00
学 历	本 科	14	35.00
	硕士生	21	52.50
	博士及以上	5	12.50
职业/专业领域	管 理	1	2.50
	法 律	3	7.50
	金 融	3	7.50
	IT/互联网	17	42.50
	文学艺术	12	30.00
	医 学	2	5.00
	新闻传媒	2	5.00

本研究使用量表工具对被试者的认知评估、情感反应和应对意图 3 个阶段中的实验观测变量进行测量,采用李克特 5 级量表收集实验数据,所有问项按照受试者的内心感受分为非常不同意、比较不同意、一般、比较同意以及非常同意,所有构念测量量表题项如表 4 所示。同时,实验中还会通过回认检测的方式对答题有效性进行验证。

4 数据分析

4.1 信效度检验与共线性评估

4.1.1 信度检验

用 SmartPLS3.0 统计软件对测量模型的测量属性进行验证性因子分析,如表 5 结果显示,测量模型的信度方面,所有变量的 Cronbach's Alpha 值均高于 0.7,组合信度(CR)均大于 0.7,表明测量模型具有良好的内部一致性和稳定性。

4.1.2 效度检验

本研究采用 SPSS 对数据进行了 KMO、Bartlett 等测试,表 6 结果显示 KMO 值为 0.84, Bartlett 球形检验中 $P<0.05$,说明适合进行结构方程模型分析,变量测量项之间具备相关性。

为判断各潜变量之间是否存在显著区别,通过平均提取方差值(AVE)计算各变量之间的相关系数

表 4 潜在变量及测量项
Tab. 4 Potential Variables and Measurement Items

变量名	题 项	来 源
AIGC 透明度	AT1 AI 搜索信息的算法机制和生成逻辑被公开发布并让我理解	Bahar M 等 ^[59]
	AT2 AI 搜索信息的评估方法和标准被公开发布并让我们理解	
	AT3 AI 搜索信息的整个过程都可以向我们解释清楚	
AIGC 可控度	AC1 我可以根据不同需求来改变 AI 搜索引擎生成的信息	Kautish P 等 ^[72]
	AC2 我感到自己可以操控并调整 AI 搜索引擎生成的信息	
	AC3 我感到自己总是可以控制 AI 生成我想要的信息	
AI 焦虑	AA1 使用 AI 服务搜索信息的过程让我感到很焦虑	Johnson D G 等 ^[73]
	AA2 学习使用这个 AI 搜索引擎让我感到很焦虑	
	AA3 浏览这个 AI 搜索返回的结果让我感到很焦虑	
AI 身份威胁	AIT1 如果我热情地使用 AI,我的同龄人就会失去对我的尊重	Mirbabaie M 等 ^[69]
	AIT2 如果我支持人工智能,会被认为是同龄人中一个可怜的成员	
	AIT3 使用 AI 让我感到不自信,无法很好地理解事情或完成工作	
	AIT4 使用人工智能让我觉得我少了一些让我脱颖而出的品质	
AIGC 信息回避	AIA1 如果继续使用,我会选择性地接收 AI 生成的信息	Dai B 等 ^[74]
	AIA2 如果继续使用,我会拒绝接收 AI 生成的一些信息	
	AIA3 如果继续使用,我会有意回避 AI 生成的一些信息	

表 5 测量模型的信度与效度指标
Tab. 5 Reliability and Validity Indicators of Measuring Model

变 量	测量项	标准化因子载荷	AVE	CR	Cronbach'α	VIF
AIGC 透明度	AT1	0.893	0.828	0.935	0.896	2.457
	AT2	0.931				3.335
	AT3	0.905				2.732
AIGC 可控度	AC1	0.939	0.885	0.959	0.953	3.086
	AC2	0.941				3.778
	AC3	0.943				4.188
AI 焦虑	AA1	0.887	0.805	0.925	0.879	2.366
	AA2	0.885				2.273
	AA3	0.919				2.842
AI 身份威胁	AIT1	0.759	0.653	0.882	0.834	1.910
	AIT2	0.77				1.962
	AIT3	0.833				1.559
	AIT4	0.864				1.898
AIGC 信息回避	AIA1	0.939	0.868	0.952	0.925	3.621
	AIA2	0.903				3.120
	AIA3	0.953				3.637

表 6 KMO 值及 Bartlett 球形检验结果
Tab. 6 KMO Values and Bartlett Sphericity Test Results

项 目		结 果
KMO 检验值		0.84
Bartlett 球形检验	近似卡方	555.454
	df	120
	P 值	0

从而进行区分效度检验。结果(参见表 7)显示,所有构念的内部相似度(粗体部分)均大于外部相似度,证明模型的构念间有着良好的区分度。

4.1.3 共线性分析

在进行 PLS-SEM 分析时,共线性评估是一个重要的质量标准。如果预测构面之间存在共线性,路径系数可能会出现偏差。通过指标方差膨胀因素

表 7 潜在因子 AVE 值的平方根与因子间的相关系数
Tab. 7 Square Root of Latent Factor AVE Value and Correlation Coefficient Between Factors

变 量	1	2	3	4	5
AIGC 透明度	0.909				
AIGC 可控度	-0.712	0.941			
AI 焦虑	-0.752	-0.651	0.897		
AI 身份威胁	-0.829	-0.631	0.763	0.808	
AIGC 信息回避	-0.681	-0.499	0.616	0.514	0.932

(VIF)评估共线性问题,一般来说,当 VIF 值在 0~10 时,表示不存在显著多重共线性;当在 10~100 时,存在多重共线性;而当 VIF 超过 100 时,则存在严重多重共线性。因此,通常要求模型的 VIF 值低于 5。本研究中各潜变量的内部 VIF 和外部 VIF 均小于 5,具体数值如表 5 和表 8 所示,可以看出本

研究各潜变量之间不存在显著多重共线性。

4.2 路径分析

采用结构方程模型来验证所提出的研究模型,结果如表 9 所示。首先,测试 AIGC 的透明度、可控度以及 AI 焦虑与 AIGC 信息回避之间的直接关系。结果表明,AIGC 透明度与 AI 焦虑呈显著负向

表 8 内部模型的 VIF 值
Tab. 8 VIF Value of Internal Model

	AT	AC	AA	AIT	AIA	AIT×AT	AIT×AC
AT			3.662				
AC	1.005		3.194				
AA							
AIT	1.008						
AIA	1.018		3.476				
AIT×AT	1.006						
AIT×AC	1.018						

相关($\beta = -0.65$, $p < 0.001$), AIGC 可控度与 AI 焦虑呈显著负向相关($\beta = -0.617$, $p < 0.001$), AIGC 透明度与 AIGC 信息回避呈显著负向相关($\beta = -0.215$, $p < 0.01$), AIGC 可控度与 AIGC 信息回避呈显著负向相关($\beta = -0.244$, $p < 0.001$), 因此, 两组自变量对 AI 焦虑和 AIGC 信息回避均有抑制作用, 随着 AIGC 透明度和 AIGC 可控度的提升, 也即 AIGC 信息不确定性的下降, 用户的 AI 焦虑会得到缓解, 并且更不倾向于采取 AIGC 信息回避行为。这些发现证实了上文提出的 H1~H4。

表 9 结构方程模型路径系数及显著性水平
Tab. 9 Path Coefficients and Significance Levels of Structural Equation Model

关系路径	路径系数	样本均值	T 值	显著性
AA→AIA	0.655***	0.646	6.063	显著
AC→AA	-0.617***	-0.615	7.747	显著
AC→AIA	-0.244***	-0.251	3.276	显著
AIT→AA	0.164	0.147	1.426	不显著
AT→AA	-0.650***	-0.644	8.453	显著
AT→AIA	-0.215***	-0.221	3.772	显著
AIT×AT→AA	-0.122*	0.027	2.082	显著
AIT×AC→AA	-0.174*	0.088	2.085	显著

注: ***表示 $p < 0.001$, *表示 $p < 0.05$ 。

本研究考察了 AI 焦虑对 AIGC 透明度与 AIGC 信息回避之间关系的中介作用, 以及 AI 焦虑对 AIGC 可控度与 AIGC 信息回避之间关系的中介作用。AI 焦虑与 AIGC 信息回避呈正向相关($\beta = 0.655$, $p < 0.001$), 用户感知到的 AI 焦虑会显著促进 AIGC 信息回避行为。研究结果支持 AI 焦虑中介关系的

假设 H5。

研究结果支持 AI 身份威胁对 AIGC 透明度与 AI 焦虑之间的负向关系强度的正向调节作用($\beta = -0.122$, $p < 0.05$), 以及对 AIGC 可控度与 AI 焦虑之间的负向关系强度的正向调节作用($\beta = -0.174$, $p < 0.05$), 即 AI 身份威胁水平会提升 AIGC 透明度、AIGC 可控度与 AI 焦虑之间的关系强度。AI 身份威胁感较高的人群里, 用户 AIGC 感知透明度和感知可控度水平对 AI 焦虑的影响更大, H6、H7 假设成立。本研究结构方程模型路径分析结果如图 4 所示。

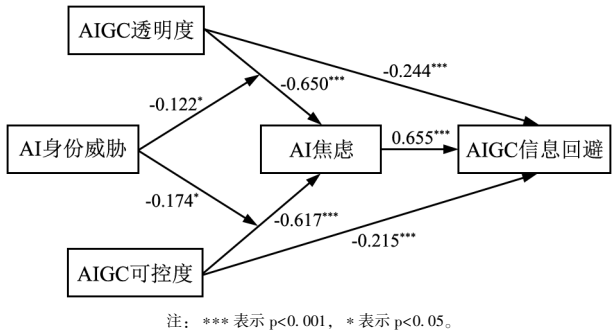


图 4 结构方程模型路径分析图

Fig. 4 Path Analysis Diagram of Structural Equation Model

为了更加直观地分析 AI 身份威胁的调节效应, 将 AIGC 透明度、AIGC 可控度以及 AI 身份威胁通过均值增减 1 个标准差实现均值漂移, 处理为高水平 and 低水平两种类别的分类变量, 目的是分析 AI 身份威胁水平高或低时 AIGC 透明度和 AIGC 可控度对 AI 焦虑的负向影响程度是否发生变化, 结果如图 5 所示。根据图 5 可知, 在使用 AIGC 应用时, 对 AI 身份威胁的评估越高, 用户越可能产生焦虑情绪, 且 AI 身份威胁对低 AIGC 透明度和低 AIGC

可控度的调节效应更强,也越可能产生 AI 焦虑情绪。受试者对 AI 身份威胁的评估较高会使其 AI 焦虑情绪随着对可控度和透明度评估水平的降低而显

著增高;而对于持有 AI 较低身份威胁观念的受试者,不确定性认知的差异对于 AI 焦虑情绪的影响程度相对较弱。

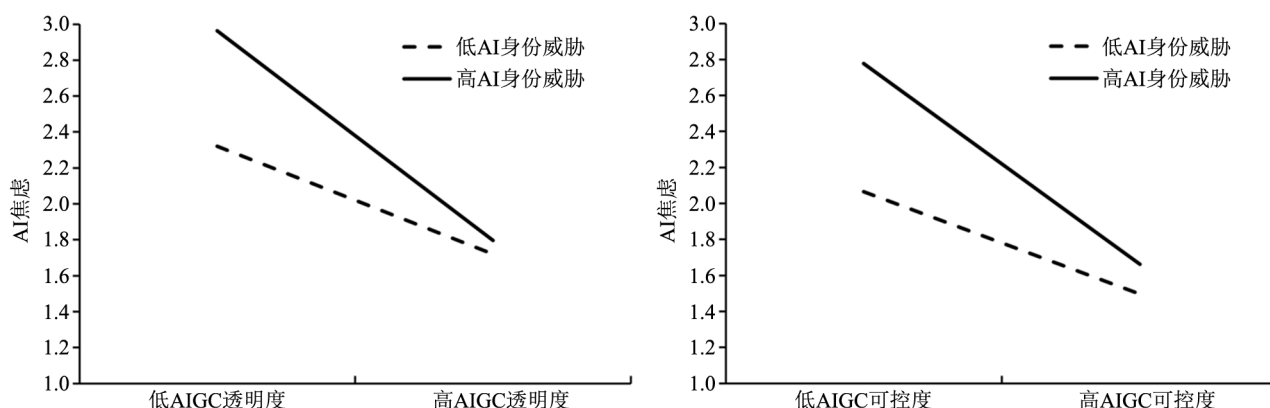


图 5 AI 身份威胁对 AIGC 透明度、AIGC 可控度与 AIGC 焦虑关系的调节作用

Fig. 5 Moderating Effect of AI Identity Threat on the Relationship Between AIGC Transparency, AIGC Controllability and AIGC Anxiety

5 研究结论与启示

5.1 研究结论

本研究基于不确定性管理理论、情绪认知理论和 C-A-C 框架构建了 AIGC 信息回避影响因素模型,探究了 AI 焦虑如何影响人们对 AIGC 的信息回避,并对 AI 身份威胁的调节作用进行了检验。

研究结果表明,AIGC 信息不确定性与 AI 焦虑之间存在显著的正相关。即,当 AIGC 信息的透明度和可控度较高时,用户对 AIGC 的信息回避意愿相对较弱,这说明信息的透明度和可控度能够缓解人们的 AI 焦虑,提高用户对 AIGC 信息的接受度。然而,当 AIGC 信息的透明度和可控度较低时,用户对 AIGC 信息的回避意愿则变得更为强烈,这可能是由于信息的缺乏和不透明会加剧人们对 AI 的担忧和恐惧。透明度帮助用户理解 AIGC 的工作原理,而可控度则让用户感觉自己生成的内容具有主导权。两者相结合,可以增强用户对 AIGC 技术的信任,减少对技术失控的恐惧,以及对可能的负面后果的担忧。

此外,AI 焦虑在 AIGC 信息不确定性与用户的回避意愿之间起到了中介作用,这说明 AI 焦虑是影响用户对 AIGC 信息接受度的重要因素。Cezar B G D S 等^[75]在大数据环境中同样验证了焦虑情绪和用户信息回避行为呈正相关。本研究扩展了信息回避的场景,并引入了 AI 焦虑作为新的焦虑情绪

类型。在 AIGC 应用或服务的情境下,AI 焦虑对信息回避有正向影响,即用户的 AI 焦虑情绪会导致他们更倾向于避免或延迟与 AIGC 相关的信息的接触和使用。这种影响可能会限制 AIGC 技术的普及和应用,因此,在设计和推广 AIGC 应用时,应充分考虑如何降低用户的焦虑情绪,提高他们对新技术的接受度。

同时,AI 身份威胁作为用户对 AI 技术可能损害个人自我信念的预期,对不确定性认知和焦虑情绪之间的关系有显著的调节作用。即 AI 身份威胁程度越强,越会加剧不确定性认知对 AI 焦虑的作用;反之,如果 AI 身份威胁程度较弱,那么即使用户对 AIGC 持有一定水平的不确定性认知,也不一定会引发太多的 AI 焦虑。以往关于身份威胁的研究表明^[44],在机器人服务的场景中,身份威胁显著调节了服务提供者类型通过社会在场对满意度的间接影响,这意味着当机器人服务员引发身份威胁时,可能降低客户的满意度。本研究将身份威胁在 AIGC 应用中的概念进行重新界定为 AI 身份威胁,并从认知情绪视角,验证了 AI 身份威胁对于 AI 焦虑的调节效应。此发现也可以解释 AIGC 应用中信息回避行为相关的矛盾现象,如在一些内容社区平台(如 B 站、小红书、Lofter 等)中 AI 绘画被极力回避,而 AI 音乐创作(如 AI 孙燕姿、AI 周杰伦)却没有被同样程度地回避。这种差异可能

源于不同领域中 AI 技术对个人身份威胁的感知程度不同。在绘画领域, AI 技术的介入可能被视为对艺术家个人创造力的直接挑战, 从而引发较强的 AI 身份威胁观念, 以致更有可能导致 AIGC 信息回避行为; 而在音乐创作领域, AI 可能更倾向于辅助工具, 帮助音乐家拓展创作的可能性, 因此用户的 AI 身份威胁程度较弱, 即使受到了不确定性认知的刺激, 相比之下用户对 AI 音乐创作的接受度相对较高, 信息回避行为也较少。这种理解有助于更好地把握用户对 AIGC 技术的态度和行为, 为 AIGC 技术的设计和推广提供指导。

5.2 研究启示

当前正处于 AIGC 技术架构体系快速变迁的阶段, AI 能力覆盖学习和执行聚焦执行与社会协作环节, AIGC 技术研发开始注重人机交互协作, 注重人类对人工智能的反馈训练。

人工智能的不断发展也引发了很多值得思考的问题: 如, 为何图片生成类 AI 应用在职业画手群体中遭遇抵制的同时, “AI 换脸” “AI 歌手” 等音视频模态的 AIGC 信息却受到广泛欢迎? 因此, 需要重视 AIGC 模态特征和信息不确定性以及其对于用户体验和行为相关的影响, 未来可以建立交互式反馈与学习机制, 让用户真实体验和了解到自己的操作如何影响 AIGC 系统的输出, 从而缓解用户的“黑箱” 焦虑, 提升用户体验。此外, 还可以赋予用户个性化定制的权利: 如, 允许用户对 AIGC 生成的内容进行个性化定制和编辑, 包括调整内容的风格、语调、结构等其他参数, 通过提升用户的参与感和控制感, 在增强用户对 AIGC 系统信任度的同时, 还能够为系统收集更多的用户反馈和数据, 从而进一步改进和优化系统的感知控制能力, 实现技术与用户的和谐共生。

不仅如此, 由于用户对 AIGC 的信息回避往往源于对其技术原理、应用场景以及潜在风险的误解或误解, 这种回避态度阻碍了公众对 AIGC 价值的认同, 限制了 AIGC 技术的广泛应用。因此, 需要专业培训与教育项目来改善这一现状, 通过制作宣传材料, 如视频、信息图表和案例研究, 详细解释 AIGC 的工作流程, 包括数据的收集、处理和内容生成的每个环节, 培养用户的数智素养, 建立技

术信任, 增强身份信念。

与此同时, 需要持续反思和探索人与技术的关系, 只有确保技术发展过程中人类的需求、情感和价值观得到充分尊重和重视, 才能让用户厘清 AI 作为“效率助手” 而非“竞争对手” 的社会角色定位, 从而可以在享受技术便利和承担技术威胁与焦虑风险之间找到平衡。

6 结 语

本文研究了用户对人工智能生成内容(AIGC)信息回避的成因与作用机制, 揭示了 AI 身份威胁在调节用户焦虑中的重要作用。这一发现有助于理解用户对 AIGC 技术的信息回避态度与对 AIGC 技术积极使用之间的矛盾现象, 也为 AIGC 相关产品功能设计和 AIGC 行业信息标准制定提供了指导意见。

实验任务设计仅局限于文本生成类 AIGC, 未来可以尝试对图片生成、音视频生成以及跨模态生成类 AIGC 使用场景进行验证, 亦可对比分析不同类别 AIGC 的用户信息回避影响因素以及内在作用机制之间的差异。此外, 由于研究样本的局限, 未来需尽可能扩大样本容量, 并注重选取各行业各领域的样本进行比较, 尽可能减少样本误差; 同时, 由于用户的不确定性认知或可能在不同人群、不同应用情景中产生不同特征, 未来需针对特定群体(如老年群体、儿童群体等)或某一类 AIGC 应用的不确定性认知特征进行深入探索, 以期丰富 AIGC 用户信息行为影响因素的理论研究。

参 考 文 献

- [1] Shao L, Chen B, Zhang Z, et al. Artificial Intelligence Generated Content(AIGC) in Medicine: A Narrative Review [J]. Mathematical Biosciences and Engineering, 2024, 21 (1): 1672-1711.
- [2] 王洁. IDC 发布 2024 年 AIGC 应用层十大趋势 [N]. 中国信息化周报, 2024-01-22, (19).
- [3] NewsGuard. Tracking AI-enabled Misinformation: 739 Unreliable AI-Generated News Websites (and Counting), Plus the Top False Narratives Generated by Artificial Intelligence Tools [EB/OL]. <https://www.newsguardtech.com/special-reports/ai-tracking-center/>, 2024-03-04.
- [4] 詹希旎, 李白杨, 孙建军. 数智融合环境下 AIGC 的场景化应用与发展机遇 [J]. 图书情报知识, 2023, 40 (1): 75-85, 55.
- [5] 储节旺, 罗怡帆. 人工智能生成内容赋能图书馆知识服务的途径研究 [J]. 情报理论与实践, 2024, 47 (8): 34-42.

- [6] 刘逸伦, 黄薇, 张晓君, 等. AIGC 赋能的科技情报智能服务: 特征、场景与框架 [J]. 现代情报, 2023, 43 (12): 88-99.
- [7] 王树义, 张庆薇, 张晋. AIGC 时代的科研工作流: 协同与 AI 赋能视角下的数字学术工具应用及其未来 [J]. 图书情报知识, 2023, 40 (5): 28-38, 126.
- [8] Panda S, Chakravarty R. Adapting Intelligent Information Services in Libraries: A Case of Smart AI Chatbots [J]. Library Hi Tech News, 2022, 39 (1): 12-15.
- [9] Chaudhry I S, Sarwary A M, El Refae G A, et al. Time to Revisit Existing Student's Performance Evaluation Approach in Higher Education Sector in a New Era of ChatGPT: A Case Study [J]. Cogent Education, 2023, 10 (1): 2210461.
- [10] 宋士杰, 赵宇翔, 朱庆华. 从 ELIZA 到 ChatGPT: 人智交互体验中的 AI 生成内容(AIGC)可信度评价 [J]. 情报资料工作, 2023, 44 (4): 35-42.
- [11] Liu Y, Wang S, Yu G. The Nudging Effect of AIGC Labeling on Users' Perceptions of Automated News: Evidence from EEG [J]. Frontiers in Psychology, 2023, 14 (12): 1277829.
- [12] Almaiah M A, Alfaisal R, Salloum S A, et al. Examining the Impact of Artificial Intelligence and Social and Computer Anxiety in E-Learning Settings: Students' Perceptions at the University Level [J]. Electronics, 2022, 11 (22): 3662.
- [13] 赵静, 倪明扬, 张倩, 等. AIGC 重构研究生学术实践: 持续使用意愿影响因素研究 [J]. 现代情报, 2024, 44 (7): 34-46.
- [14] Suseno Y, Chang C, Hudik 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, 2022, 33 (6): 1209-1236.
- [15] Xenidis R. Tuning EU Equality Law to Algorithmic Discrimination: Three Pathways to Resilience [J]. Journal of European and Comparative Law, 2020, 27 (6): 736-758.
- [16] Commerford B P, Dennis S A, Joe J R, et al. Man Versus Machine: Complex Estimates and Auditor Reliance on Artificial Intelligence [J]. Journal of Accounting Research, 2022, 60 (1): 171-201.
- [17] Cheng X, Zhang X, Cohen J, et al. Human vs. AI: Understanding the Impact of Anthropomorphism on Consumer Response to Chatbots from the Perspective of Trust and Relationship Norms [J]. Information Processing and Management, 2022, 59 (3): 102940.
- [18] Xie Z, Yu Y, Zhang J, et al. The Searching Artificial Intelligence: Consumers Show Less Aversion to Algorithm-Recommended Search Product [J]. Psychology Marketing, 2022, 39 (10): 1902-1919.
- [19] Polyportis A. A Longitudinal Study on Artificial Intelligence Adoption: Understanding the Drivers of ChatGPT Usage Behavior Change in Higher Education [J]. Frontiers in Artificial Intelligence, 2024, 6 (1): 1324398.
- [20] Brashers D E, Goldsmith D J, Hsieh E. Information Seeking and Avoiding in Health Contexts [J]. Human Communication Research, 2002, 28 (2): 258-271.
- [21] Dali K. The Lifeways we Avoid: the Role of Information Avoidance in Discrimination Against People with Disabilities [J]. Journal of Documentation, 2018, 74 (6): 1258-1273.
- [22] 张玥, 赵恬, 黄冰冰, 等. 媒介依赖视角下社交信息回避策略对移动阅读效果的实验研究 [J]. 情报资料工作, 2024, 45 (2): 75-85.
- [23] Galarce M E, Ramanadhan S, Weeks J, et al. Class, Race, Ethnicity and Information Needs in Post-Treatment Cancer Patients [J]. Patient Education and Counseling, 2011, 85 (3): 432-439.
- [24] 巩洪村, 邓三鸿, 曹高辉. 科研知识交流情景下信息规避行为的潜在成因分析 [J]. 情报杂志, 2022, 41 (6): 189-199, 141.
- [25] Guo Y Y, Lu Z Z, Kuang H B, et al. Information Avoidance Behavior on Social Network Sites: Information Irrelevance, Overload, and the Moderating Role of Time Pressure [J]. International Journal of Information Management, 2020, 52 (6): 102067.
- [26] 陈烨, 王阳, 岳欣, 等. 运动损伤人群健康信息规避行为作用机理研究 [J]. 现代情报, 2023, 43 (1): 55-68.
- [27] 方睿. 新冠病毒疫情信息回避行为的影响因素研究 [D]. 济南: 山东大学, 2024.
- [28] 顾东晓, 孙佳月, 丁庆秀, 等. 健康信息焦虑对健康信息规避行为的影响路径研究——基于扎根理论的探索 [J]. 图书情报工作, 2023, 67 (19): 111-120.
- [29] Filieri R, McLeay F, Tsui B, et al. Consumer Perceptions of Information Helpfulness and Determinants of Purchase Intention in Online Consumer Reviews of Services [J]. Information & Management, 2018, 55 (8): 956-970.
- [30] 张皓翔. 社会化问答情境下网络用户信息规避行为的影响机制研究——基于扎根理论的探索性分析 [J]. 情报科学, 2022, 40 (12): 63-72.
- [31] 马海云, 薛翔. 用户信息搜寻到信息规避的演化机制研究——以突发公共卫生事件领域为例 [J]. 现代情报, 2024, 44 (9): 107-118, 130.
- [32] 周彤, 李玉玲, 王雪聪, 等. 大学生健康信息回避现状及影响因素研究 [J]. 图书馆学研究, 2023, (5): 70-78, 69.
- [33] 孙晓宁, 景雨田. 弹幕用户信息规避行为过程与情感演化研究 [J]. 图书情报知识, 2024, 41 (2): 87-99.
- [34] 刘百灵, 杨世龙, 李延晖. 隐私偏好设置与隐私反馈对移动商务用户行为意愿影响及交互作用的实证研究 [J]. 中国管理科学, 2018, 26 (8): 164-178.
- [35] Fishbein M, Ajzen I. Belief, Attitude, Intention and Behaviour: An Introduction to Theory and Research [J]. Philosophy & Rhetoric, 1977, 41 (4): 842-844.

- [36] Lazarus S R. The Psychology of Stress and Coping [J]. Issues in Mental Health Nursing, 2010, 7 (1-4): 399-418.
- [37] Haley W E, Levine E G, Brown S L, et al. Stress, Appraisal, Coping, and Social Support as Predictors of Adaptational Outcome Among Dementia Caregivers [J]. Psychology and Aging, 1987, 2 (4): 323-330.
- [38] Li J, Hudson S, So K K F. Hedonic Consumption Pathway vs. Acquisition-Transaction Utility Pathway: An Empirical Comparison of Airbnb and Hotels [J]. International Journal of Hospitality Management, 2021, 94: 102844.
- [39] Cao Y Y, Qin X H, Li J J, et al. Exploring Seniors' Continuance Intention to Use Mobile Social Network Sites in China: A Cognitive-Affective-Conative Model [J]. Universal Access in the Information Society, 2022, 21 (1): 71-92.
- [40] Huang D, Chen Q, Huang J, et al. Customer-Robot Interactions: Understanding Customer Experience with Service Robots [J]. International Journal of Hospitality Management, 2021, 99 (1): 103078.
- [41] Brashers E D. Communication and Uncertainty Management [J]. Journal of Communication, 2001, 51 (3): 477-497.
- [42] Yang J Z, Aloe M A, Feeley H T. Risk Information Seeking and Processing Model: A Meta-Analysis [J]. Journal of Communication, 2014, 64 (1): 20-41.
- [43] 艾文华, 胡广伟, 赵宇翔, 等. 健康信息规避行为影响因素研究: 基于元分析的探索 [J]. 情报资料工作, 2021, 42 (6): 63-73.
- [44] Lazarus R S, Folkman S. Stress, Appraisal, and Coping [M]. London: Springer New York, 1984.
- [45] Milyavsky M, Webber D, Fernandez J R, et al. To Reappraise or not to Reappraise? Emotion Regulation Choice and Cognitive Energetics [J]. Emotion, 2019, 19 (6): 964-981.
- [46] Luo M M, Chea S. Cognitive Appraisal of Incident Handling, Affects, and Post-Adoption Behaviors: A Test of Affective Events Theory [J]. International Journal of Information Management, 2018, 40 (2): 120-131.
- [47] 代宝, 杨泽国. 社交媒体用户信息回避行为的影响因素分析 [J]. 信息资源管理学报, 2022, 12 (2): 13-24.
- [48] 孙海霞. 国外健康信息规避行为研究综述 [J]. 图书情报工作, 2021, 65 (9): 138-150.
- [49] Wu L, Long A, Hu C, et al. An Identity Threat Perspective on Why and When Employee Voice Brings Abusive Supervision [J]. Frontiers in Psychology, 2023, 14 (2): 1133480.
- [50] Singh S, Olson E D, Tsai C H K. Use of Service Robots in an Event Setting: Understanding the Role of Social Presence, Eeriness, and Identity Threat [J]. Journal of Hospitality and Tourism Management, 2021, 49 (4): 528-537.
- [51] Christy Ds. The Danger Is not Machines Becoming Humans, But Humans Becoming Machines [EB/OL]. <https://news.harvard.edu/gazette/story/2023/07/Making-Algorithm-Used-In-Ai-More-Human-Like/>, 2023-10-12.
- [52] Big Think. Making Algorithm Used in AI More Human-Like [EB/OL]. <https://Bigthink.com/articles/The-Danger-Is-Not-Machines-Becoming-Humans-But-Humans-Becoming-Machins/>, 2023-12-13.
- [53] 胡丽霞, 闵庆飞, 李梦一. 电商直播技术社会示范性对消费者平台使用意向的影响 [J]. 管理科学, 2023, 36 (1): 1-15.
- [54] Venkatesh V, Thong L Y J, Chan Y K F, et al. Managing Citizens' Uncertainty in E-Government Services: The Mediating and Moderating Roles of Transparency and Trust [J]. Information Systems Research, 2016, 27 (1): 87-111.
- [55] Zimbres T M, Bell R A, Miller L M S, et al. When Media Health Stories Conflict: Test of the Contradictory Health Information Processing(CHIP) Model [J]. Journal of Health Communication, 2021, 26 (7): 460-472.
- [56] Nguyen M H, Bol N, Lustria M L A. Perceived Active Control Over Online Health Information: Underlying Mechanisms of Mode Tailoring Effects on Website Attitude and Information Recall [J]. Journal of Health Communication, 2020, 25 (4): 271-282.
- [57] Zhu K. Information Transparency of Business-to-Business Electronic Markets: A Game-Theoretic Analysis [J]. Management Science, 2004, 50 (5): 670-685.
- [58] Agozie D Q, Kaya T. Discerning the Effect of Privacy Information Transparency on Privacy Fatigue in E-Government [J]. Government Information Quarterly, 2021, 38 (4): 101601.
- [59] Bahar M, Tenzin D. Fairness, Accountability, Transparency, and Ethics(FATE) in Artificial Intelligence(AI) and Higher Education: A Systematic Review [J]. Computers and Education: Artificial Intelligence, 2023, 5 (1): 100152.
- [60] Marken R S, Mansell W. Perceptual Control as a Unifying Concept in Psychology [J]. Review of General Psychology, 2013, 17 (2): 190-195.
- [61] Elie-Dit-Cosaque M C, Pallud J, Kalika M. The Influence of Individual, Contextual, and Social Factors on Perceived Behavioral Control of Information Technology: A Field Theory Approach [J]. Journal of Management Information Systems, 2012, 28 (3): 201-234.
- [62] Li M L, Yin D X, Qiu H L, et al. Examining the Effects of AI Contactless Services on Customer Psychological Safety, Perceived Value, and Hospitality Service Quality During the COVID-19 Pandemic [J]. Journal of Hospitality Marketing Management, 2022, 31 (1): 24-48.
- [63] Li J, Huang J S. Dimensions of Artificial Intelligence Anxiety

- Based on the Integrated Fear Acquisition Theory [J]. Technology in Society, 2021, 63 (4): 101410.
- [64] Dawid B, Brenda C S. Barriers to Adopting Automated Organisational Decision-Making Through the Use of Artificial Intelligence [J]. Management Research Review, 2024, 47 (1): 64-85.
- [65] Schweitzer F, Belk R, Jordan W, et al. Servant, Friend or Master? The Relationships Users Build with Voice-Controlled Smart Devices [J]. Journal of Marketing Management, 2019, 35 (7-8): 693-715.
- [66] Park C S, Kaye B K. Smartphone and Self-Extension: Functionally, Anthropomorphically, and Ontologically Extending Self via the Smartphone [J]. Mobile Media & Communication, 2019, 7 (2): 215-231.
- [67] Petriglieri L J. Under Threat: Responses to and the Consequences of Threats to Individuals Identities [J]. Academy of Management Review, 2011, 36 (4): 641-662.
- [68] Nach H, Lejeune A. Coping with Information Technology Challenges to Identity: A Theoretical Framework [J]. Computers in Human Behavior, 2010, 26 (4): 618-629.
- [69] Mirbabaie M, Stieglitz S, Brünker F, et al. Understanding Collaboration with Virtual Assistants: The Role of Social Identity and the Extended Self [J]. Business & Information Systems Engineer-
- ing, 2020, 63 (1): 21-37.
- [70] Precedence Research. Artificial Intelligence(AI) Market [EB/OL]. <https://www.precedence-research.com/artificial-intelligence-market/>, 2024-01-06.
- [71] Questmobile. 2024 生成式 AI 及 AIGC 应用洞察报告: 头部 App 应用去重月活用户突破 5000 万, C 端、B 端机会蜂拥而至 [EB/OL]. <https://www.questmobile.com.cn/research/report/1767395734913650690/>, 2024-03-12.
- [72] Kautish P, Khare A. Investigating the Moderating Role of AI-Enabled Services on Flow and Awe Experience [J]. International Journal of Information Management, 2022, 66 (5): 102519.
- [73] Johnson D G, Verdicchio M. AI Anxiety [J]. Journal of the Association for Information Science and Technology, 2017, 68 (9): 2267-2270.
- [74] Dai B, Ali A, Wang H. Exploring Information Avoidance Intention of Social Media Users: A Cognition-Affect-Conation Perspective [J]. Internet Research, 2020, 30 (5): 1455-1478.
- [75] Cezar B G D S, Maçada A C G. Cognitive Overload, Anxiety, Cognitive Fatigue, Avoidance Behavior and Data Literacy in Big Data Environments [J]. Information Processing & Management, 2023, 60 (6): 103482.

(责任编辑: 郭沫含)

(上接第 11 页)

- [38] Dagdelen J, Dunn A, Lee S, et al. Structured Information Extraction from Scientific Text with Large Language Models [J]. Nature Communications, 2024, 15 (1): 1418.
- [39] Liu J, Wang J, Huang H, et al. Improving LLM-Based Health Information Extraction with In-Context Learning [C] //China Health Information Processing Conference. Singapore: Springer Nature Singapore, 2023: 49-59.
- [40] Jung S J, Kim H, Jang K S. LLM Based Biological Named Entity Recognition from Scientific Literature [C] //2024 IEEE International Conference on Big Data and Smart Computing (BigComp). IEEE, 2024: 433-435.
- [41] Kumar M P, Packer B, Koller D, et al. Self-Paced Learning for Latent Variable Models [C] //Proceedings of the 24th International Conference on Neural Information Processing Systems, 2010: 1189-1197.
- [42] 张宇, 刘波. 基于自步学习策略的归纳式迁移学习模型研究 [J]. 广东工业大学学报, 2023, 40 (4): 31-36.
- [43] 江雨燕, 陶承凤, 李平. 数据增强和自适应自步学习的深度子空间聚类算法 [J]. 计算机工程, 2023, 49 (8): 96-103, 110.
- [44] Huang Z H, Xu W, Yu K. Bidirectional LSTM-CRF Models for Sequence Tagging [EB/OL]. <https://arxiv.org/abs/1508.01991>, 2015-08-09.
- [45] Souza F, Nogueira R, Lotufo R. Portuguese Named Entity Recognition using BERT-CRF [EB/OL]. <https://arxiv.org/abs/1909.10649>, 2020-02-27.
- [46] Lopez P, Du C F, Cohoon J, et al. Mining Software Entities in Scientific Literature: Document-level NEW for an Extremely Imbalance and Large-scale Task [C] //Proceedings of the 30th ACM International Conference on Information & Knowledge Management, 2021: 3986-3995.
- [47] Dai Z J, Wang X T, Ni P, et al. Named Entity Recognition Using BERT BiLSTM CRF for Chinese Electronic Health Records [C] //2019 12th International Congress on Image and Signal Processing, BioMedical Engineering and Informatics (CISP - BMEI). IEEE, 2019: 1-5.
- [48] Zhang H, Zhang C Z, Wang Y Z. Revealing the Technology Development of Natural Language Processing: A Scientific Entity-Centric Perspective [J]. Information Processing & Management, 2024, 61 (1): 103574.

(责任编辑: 郭沫含)