

基于 C-A-C 的生成式 AI 用户间歇性中辍行为研究

周涛¹ 张春雷¹ 邓胜利²

(1. 杭州电子科技大学管理学院, 浙江 杭州 310018; 2. 武汉大学信息管理学院, 湖北 武汉 430072)

摘要: [目的/意义] 用户的间歇性中辍作为一种消极行为, 将影响生成式 AI 的用户保持及获取持续竞争优势。因此, 有必要研究用户间歇性中辍的形成机理, 发现显著的影响因素。[方法/过程] 基于认知—情感—意愿(Cognition-Affect-Conation, C-A-C), 从“使能”与“抑制”双重视角研究了生成式 AI 用户间歇性中辍行为。使能因素包括隐私担忧、信息幻觉、认知失调; 抑制因素包括智能化、拟人化、个性化、情感承诺。采用结构方程模型(Structural Equation Modeling, SEM)和模糊集定性比较分析(fuzzy-set Qualitative Comparative Analysis, fsQCA)进行数据分析。[结果/结论] 隐私担忧和信息幻觉影响认知失调, 进而导致间歇性中辍行为。智能化、拟人化、个性化影响情感承诺, 进而对间歇性中辍产生抑制作用。研究结果显示, 生成式 AI 一方面需要缓解用户的隐私担忧, 减少信息幻觉, 从而降低用户认知失调; 另一方面, 需要通过提高系统的智能化、拟人化、个性化等功能水平, 增进用户情感承诺, 从而抑制其间歇性中辍行为。

关键词: 生成式 AI; 间歇性中辍; C-A-C; 认知失调; 情感承诺

DOI: 10.3969/j.issn.1008-0821.2025.03.004

[中图分类号] G252.0 [文献标识码] A [文章编号] 1008-0821 (2025) 03-0040-11

Research on Generative AI Users' Intermittent Discontinuance Based on the C-A-C

Zhou Tao¹ Zhang Chunlei¹ Deng Shengli²

(1. School of Management, Hangzhou Dianzi University, Hangzhou 310018, China;

2. School of Information Management, Wuhan University, Wuhan 430072, China)

Abstract: [Purpose/Significance] As a passive behavior, users' intermittent discontinuance may make it difficult for generative AI to retain users and achieve a competitive advantage. Thus, it is necessary to examine the formation mechanism of user intermittent discontinuance and find the significant factors. [Method/Process] From a C-A-C perspective, this research examined the enablers and inhibitors of generative AI users' intermittent discontinuance. Enablers include privacy concern, information hallucination and cognitive dissonance, while inhibitors include intelligence, anthropomorphism, personalization and emotional commitment. SEM and fsQCA were adopted to conduct data analysis. [Results/Conclusions] Privacy concern and information hallucination affect cognitive dissonance, which further leads to intermittent discontinuation. Intelligence, anthropomorphism, and personalization affect emotional commitment, which prevents intermittent discontinuance. The results imply that generative AI needs to mitigate users' privacy concern and reduce information hallucination in order to lower cognitive dissonance. On the other hand, generative AI needs to enhance the intelligence, anthropomorphism, and personalization in order to increase users' emotional commitment and prevent their intermittent discontinuance.

Key words: generative AI; intermittent discontinuance; C-A-C; cognitive dissonance; emotional commitment

人工智能(Artificial Intelligence, AI)技术的快速发展正在深刻地重塑人们的思维方式与认知结

构, 并持续推动传统产业格局的革新转型。2022 年, ChatGPT 的爆发性流行让生成式 AI 成为了公众关

收稿日期: 2024-05-12

基金项目: 国家社会科学基金一般项目“基于隐私悖论的生成式 AI 用户信息披露形成机理与引导策略研究”(项目编号: 24BGL310)。

作者简介: 周涛(1979-), 男, 教授, 博士, 博士生导师, 研究方向: 信息系统与电子商务。张春雷(1997-), 男, 硕士研究生, 研究方向: 电子商务。邓胜利(1979-), 男, 教授, 博士, 博士生导师, 研究方向: 用户行为与服务。

注焦点。在此背景下，国内外涌现了众多各具特色的生成式 AI 产品，如文心一言、通义千问等。这类 AI 的核心原理在于运用算法、模型以及规则从数据中学习对象特征，并基于此创造与原始资料相似却又新颖的产品或内容^[1]。不同于以往专注于信息提炼和有限预测的“分析型 AI”或“决策型 AI”，生成式 AI 具备创新性地生成与训练样本显著不同的全新内容的能力^[2]。生成式 AI 的应用场景日益丰富和多元化，并已在内容创作、个性化推荐、虚拟助手和信息检索等领域展现出广泛的应用潜力与价值。

但与此同时，随着生成式 AI 深入日常生活应用，用户对虚假信息、隐私风险等方面的担忧也日渐增长，并导致其消极行为，如间歇性中辍，即用户在使用平台的过程中出现周期性或非连续性的使用减少直至暂时弃用的现象^[3]。间歇性中辍可能导致用户流失，影响生成式 AI 获取持续竞争优势^[4]。现有研究主要关注生成式 AI 用户的积极行为，如采纳和持续使用，较少关注消极行为，如生成式 AI 用户间歇性中辍行为。基于此，本文将从认知—情感—意愿 (Cognition-Affect-Conation, C-A-C) 视角，考察影响生成式 AI 用户间歇性中辍的“使能”

与“抑制”因素。使能因素包括隐私担忧、信息幻觉、认知失调；抑制因素包括智能化、拟人化、个性化、情感承诺。研究结果将揭示生成式 AI 用户间歇性中辍行为的成因，并为生成式 AI 提供决策借鉴和参考，从而采取有效措施来抑制用户间歇性中辍，实现用户保持，获取竞争优势。

1 文献综述

1.1 生成式 AI 用户行为

生成式 AI 的快速发展与广泛应用引发了学术界对用户行为的关注。作为新兴技术，生成式 AI 用户的采纳和使用是确保其成功的首要一步，这也是已有研究的重心。表 1 列出了关于生成式 AI 用户行为的代表性文献。从表 2 可以发现，已有研究主要关注用户对生成式 AI 的采纳和持续使用，相应的理论基础包括技术采纳与使用整合模型 (UTAUT)、刺激—组织—响应 (SOR)、任务技术匹配 (TTF)、人工智能设备使用采纳 (AIDUA) 模型等，发现了绩效期望、努力期望、信息质量等因素对用户行为的作用，但较少关注用户的消极行为，如间歇性中辍。因此，有必要研究生成式 AI 用户间歇性中辍的成因，从而采取措施抑制其中辍行为，实现用户保持。

表 1 生成式 AI 用户行为研究
Tab. 1 Generative AI User Behavior Research

文献来源	理论基础	研究问题	研究结果
Menon D 等 ^[5]	UTAUT	用户对 ChatGPT 的使用	隐私关注、绩效期望、努力期望等影响使用意愿
Li W Y ^[6]	UTAUT	设计师使用 AIGC 进行辅助设计	感知焦虑、感知风险、绩效期望等影响设计师使用意愿
Faruk L I D 等 ^[7]	心理—技术视角	大学生对 ChatGPT 的采纳	感知有用性、人性化、新奇感和用户的开放性是影响大学生的采纳意向
Ma X Y 等 ^[8]	AIDUA	用户对 ChatGPT 的采纳	认知态度和情感态度影响用户的采纳意向
毛太田等 ^[9]	扎根理论	AIGC 用户采纳	信息质量、用户感知和技术特性影响 AIGC 采纳意向
张海等 ^[10]	扎根理论	ChatGPT 用户使用意愿	主体因素、技术因素、信息因素和社会环境等影响用户使用意愿
Pham H C 等 ^[11]	SOR	游客对 ChatGPT 的持续使用	感知温暖、感知能力等影响游客的持续使用
Makkonen M 等 ^[12]	TTF	用户持续使用生成式 AI 的意愿	职业特性与用户属性影响感知有用性及用户持续意愿
赵静等 ^[13]	UTAUT2 与 TTF	研究生持续使用 AIGC 的意愿	绩效期望、努力期望、个体创新性等影响持续使用意愿

1.2 间歇性中辍

中辍反映了一种采纳后行为,是指用户在先前采用了某种创新并且使用了一段时间之后,又决定拒绝或中止使用该创新的一种消极使用行为^[14]。随着对该现象的深入研究,有学者发现用户的决策并非一成不变的,关注到用户在中止使用某项技术后仍会重新采纳的现象^[15]。因此,研究人员将“中辍”概念延伸为“间歇性中辍”,即用户在采纳某种创新后,在一段时间内中止使用或降低使用频率,但后来又再次使用或恢复正常使用频率,并可能反复循环这个过程^[3]。由此可见,间歇性中辍行为并不意味着彻底地放弃使用,而是一种动态的“中止—恢复”的过程。

已有文献考察了社交网络、短视频、智能健康硬件等背景下的间歇性中辍行为,表 2 列出了一些代表性文献。这些文献采用的理论基础包括社会比较理论、理性行为理论、SOR、压力—应对—结果(SSO)模型等,发现了负面情感,如倦怠、疲劳、不满意等对间歇性中辍的作用。从表 2 可以发现,已有文献聚焦于社交媒体用户间歇性中辍行为,较少考察作为新兴应用的生成式 AI 用户的间歇性中辍行为。此外,已有文献主要关注负面情感对间歇性中辍行为的“使能”影响,较少考察“抑制”间歇性中辍行为的因素。基于此,本文将综合考察使能因素和抑制因素对生成式 AI 用户间歇性中辍行为的影响。

表 2 间歇性中辍行为研究
Tab. 2 Intermittent Discontinuance Behavior Research

文献来源	理论基础	研究问题	研究结果
甘春梅等 ^[16]	社会比较理论	社交网络用户间歇性中辍行为	倦怠和嫉妒影响用户间歇性中辍行为
张敏等 ^[17]	SOR	社交媒体用户中辍行为	功能过载、网络外部性等影响用户负面情感,进而导致间歇性中辍
Zhang S W 等 ^[18]	SSO	社交网络用户中辍行为	功能过载、信息过载和社交过载通过疲劳和不满意影响中辍行为
李君君等 ^[19]	—	微博用户间歇性中辍行为特征	微博平台的环境、内容等促使中辍行为产生
杜诗瑶 ^[20]	扎根理论	短视频平台用户中辍行为	平台建设、内容质量等影响用户使用体验和 中辍行为
甘春梅等 ^[21]	SOR	短视频用户间歇性中辍意愿	倦怠和心流体验影响间歇性中辍意愿
王文韬等 ^[22]	理性行为理论	短视频用户间歇性中辍行为	感知成本、系统质量、群体规范影响用户中 辍行为
沈校亮等 ^[23]	矛盾态度视角	智能健康硬件用户间歇性中止行为	用户的矛盾态度和情绪效价波动影响间歇性 中止行为

2 研究模型和假设

2.1 C-A-C

认知—情感—意愿(C-A-C)认为个体行为决策包括认知、情感和意愿三大构成要素^[24]。认知体现在对事物本质的理解与解析,比如解读信息的实际含义;情感则是在认知基础上对事物所产生的主观情绪反应,如对信息的好恶感受;意愿则作为桥梁连接认知、情感与实际行动,表征行为的潜在动机,如行为意向。C-A-C 在信息系统用户行为研究中得到了广泛应用,包括对健康类应用程序的接

受^[25]、消费者对 VR 的购买意愿^[26]、社交媒体用户信息规避^[24]、社交媒体用户倦怠行为^[27]等。C-A-C 为了解用户的行为决策过程提供了一个有用的视角。因此,本文将从 C-A-C 视角考察生成式 AI 用户间歇性中辍行为,研究结果将揭示生成式 AI 用户间歇性中辍行为的内在形成机理。

基于 C-A-C,本文从“使能因素”与“抑制因素”视角考察了生成式 AI 用户间歇性中辍行为。使能因素包括隐私担忧、信息幻觉、认知失调,抑制因素包括智能化、拟人化、个性化、情感承诺。

其中, 认知失调和情感承诺是情感因素, 其他因素是认知因素。隐私担忧反映了用户对个人隐私数据可能被不恰当地使用或泄露的忧虑, 信息幻觉指生成式 AI 可能提供虚构、误导性信息。智能化、拟人化、个性化分别反映了用户对生成式 AI 技术专业性的、类人特征水平以及个性化需求满足能力的认知评价。认知失调反映了用户对生成式 AI 的期望和实际感受之间的差距, 情感承诺反映了用户与生成式 AI 的情感联系, 如依恋、认同等。研究模型如图 1 所示。

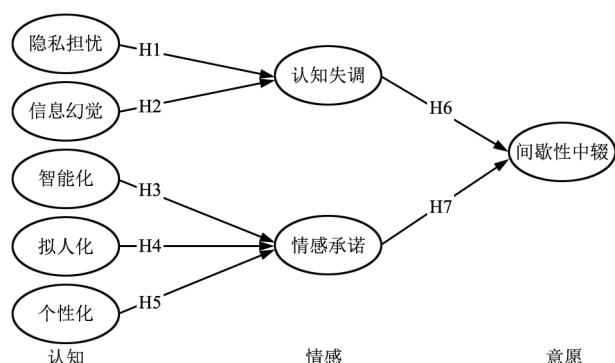


图 1 研究模型
Fig. 1 Research Model

2.2 隐私担忧

隐私担忧是指个人在披露个人数据时对其信息保密性的担忧^[28]。生成式 AI 通常需要大量数据进行训练以优化其性能和提高准确性, 这些数据可能包括用户的个人资料、搜索历史、对话记录等隐私信息。大模型内部工作原理对于普通用户来说相对复杂且不透明, 用户不清楚他们的数据如何被处理以及是否有可能被泄露。此外, 生成式 AI 模型具有强大的上下文理解和记忆能力, 可能在后续交互中“回忆”起以前的对话内容, 这可能导致存储在模型中的个人信息在未经授权的情况下被复现或滥用。

隐私担忧表明用户在享受 AI 带来的便利性的同时, 可能会牺牲个人信息隐私。因此, 用户内心产生了一种认知冲突: 他们既希望利用 AI 的强大功能, 又担心这样做将引发隐私风险, 即“隐私悖论”。这种内在的冲突加剧了用户的认知失调程度, 促使他们在继续使用 AI 与保护个人隐私之间权衡, 可能导致一系列的行为调整, 如减少使用频率、寻求更安全的替代方案或重新评估对 AI 的信任度。陈昊等^[29]发现, 社交媒体用户的隐私关注会引发

负向情感响应, 进而影响其持续使用意愿, Menon D 等^[5]认为隐私关注影响用户对 ChatGPT 的使用。因此, 本文提出以下假设:

H1: 隐私担忧正向影响用户的认知失调

2.3 信息幻觉

信息幻觉是指大模型在未基于真实数据的情况下, 通过对其训练数据的不合理外推或创造性解释所生成的低质量内容, 包括看似合理实则无意义、超现实甚至幻想性的图像、文本、声音或视频内容^[30]。如果生成式 AI 提供了虚假或不真实信息, 用户可能会基于错误的信息形成决策或信念, 当发现真相时, 其行为或信念与新的信息之间就可能产生矛盾, 导致用户感到不安、困惑或焦虑, 从而产生认知失调^[31]。前后相互矛盾或含糊不清的信息会导致用户在理解和判断上出现困扰, 尤其是在面临重要决策时, 这种情况会增加认知失调的可能性^[32]。若用户发现生成式 AI 存在信息幻觉, 其对所获取信息的信任感也会下降, 这与其对 AIGC 的预期相悖, 进一步加重认知失调。张皓翔^[33]发现, 在社会化问答社区中, 信息质量和信息需求引发信息与需求不匹配的冲突, 导致用户的认知失调和信息规避行为。因此, 本文提出以下假设:

H2: 信息幻觉正向影响用户的认知失调

2.4 智能化

智能化作为人工智能的核心特征, 指 AI 系统所展现的专业技能、知识底蕴等。Moussawi S 等^[34]将智能化定义为用户感知到的人工智能行为的有效性、自主性和自然语言处理与生成能力。当生成式 AI 能够理解用户需求、自主响应并高效完成任务时, 用户可能认为其具备智能特性并产生较高程度的信任。在实际应用中, 智能化的生成式 AI 通过自动处理大量数据和复杂任务, 显著提升了工作效率, 减轻了用户的负担。这种便捷性和效率的提升往往促使用户对其产生依赖, 并逐步深化情感连接。研究表明, 用户在与 AI 互动的过程中, 能够建立起类似与人类交往般的亲密感和积极热情, 进而促进用户的承诺和长期使用行为^[35]。

随着机器学习技术的进步, 智能化 AI 还展现出自我学习和持续优化的能力, 意味着其服务质量和适应性会随时间不断提高, 从而让用户感到自身

需求得到了充分理解和回应,进一步强化了情感联系。Rafiq F 等^[36]的研究显示, AI 聊天机器人的智能性影响用户认知态度和情感态度,进而影响采纳意向。Moussawi S 等^[37]的研究发现,智能性影响用户对个人智能代理的有用性、易用性感知和初始信任。高度智能化 AI 将超越用户预期,创造出新颖的内容和灵感火花,其创新能力带来的惊喜体验会进一步激发用户的情感投入与长期承诺。因此,本文提出以下假设:

H3: 智能化正向影响用户的情感承诺

2.5 拟人化

拟人化指的是赋予机器人等非生物实体以人类的特征表现,如仿人类的外观和情感行为模拟^[38]。拟人化设计在生成式 AI 中扮演了关键角色,它通过赋予 AI 系统更加人性化的特征和交互方式来满足用户的社会需求,增强了用户与技术之间的情感纽带。当 AI 能够模拟人类的对话习惯、情绪反应、个性特质时,用户更易产生如同与真人交流的真实体验感,并由此激发共情反应,进而可能对 AI 系统形成情感认同和信任^[39]。已有研究表明,声控 AI 设备的拟人化程度影响用户对其产生的情感依恋程度^[40]。Zhang A D 等^[41]发现,在智能家居情境下,拟人化水平以及社会角色设定影响用户的情感依恋、信息披露倾向及满意度。此外,拟人化产品由于类人特征带来的亲切感和熟悉感,使得它们能被用户自然接纳并建立信任关系,在互动过程中往往能够引发积极的情绪反应,如愉悦感和满足感,从而刺激用户的探索欲望和兴趣,进一步提升用户对 AI 的忠诚度。Moussawi S 等^[37]也证实了个人智能代理的拟人化特性对提升用户感知愉悦具有显著作用。因此,本文提出以下假设:

H4: 拟人化正向影响用户的情感承诺

2.6 个性化

个性化指的是基于用户的独特需求与偏好而定制服务的过程^[42],对于满足用户在自我表达、效率提升及所有权感知等层面的诉求至关重要。生成式 AI 通过学习和理解用户的偏好、行为模式及情感反应,能够提供高度精准且定制化的体验和服务。例如,生成式 AI 平台 Midjourney 可以根据用户提供的文本描述创建具有个人特色的艺术作品或图像,

此类深度个性化的互动体验会增加用户对平台的认同感和依赖度,提升情感承诺。Choi N^[43]发现,个性化对建立用户依恋具有显著作用。此外,Sheng H 等^[44]也证实,个性化设计对用户形成情感依恋以及功能依恋具有显著作用。当用户感知到生成式 AI 的独特价值和个性化关怀时,会更倾向于将其融入日常生活,并由此建立起深厚的情感纽带,进而发展为强烈的情感承诺。因此,本文提出以下假设:

H5: 个性化正向影响用户的情感承诺

2.7 认知失调

认知失调理论认为,个体在意识到自身态度内在的不一致性或行为与态度间的矛盾关系时,会本能地体验到心理不适,并倾向于采取措施来缓解这种冲突状态^[45]。生成式 AI 用户可能一方面认可并享受 AI 带来的便利和效率提升,另一方面又担忧其个人数据的安全性和隐私权受到侵犯。此外,生成式 AI 有时会输出看似合理实则虚假或误导性的内容,期望与实际体验之间的差距导致用户对其功能性和可靠性产生怀疑,这种内在的不一致性导致了认知失调状态。面对认知失调的压力,用户可能采取缓解策略,其中之一就是选择间歇性地中止或减少使用生成式 AI^[46]。这种间歇性中辍可以理解为用户试图通过暂时远离 AI 来减轻由于认知冲突所带来的压力,并可能在此期间寻求自我价值确认、重构对 AI 技术的认知或调整自己的使用方式。已有研究发现,认知失调会导致社会化商务用户产生潜水、抱怨和间断使用等消极使用行为^[31],诱发社会化问答社区用户信息规避行为^[33],影响短视频社交媒体用户不持续使用意向^[47]。因此,本文提出以下假设:

H6: 认知失调正向影响用户的间歇性中辍行为

2.8 情感承诺

承诺理论在组织行为学中指的是组织成员持续忠诚于组织的意愿^[48]。在生成式 AI 情境下,情感承诺反映了用户维持关系的意愿,包括对生成式 AI 的认同和依恋^[49]。例如,自媒体创作者可能会运用生成式 AI 来构建文章结构或启发创新思维。当他们观察到此类工具能够实质性地提升创作效率,并助力产出更多优质内容时,他们会认同生成式 AI 的价值,并形成情感上的依赖与承诺。情感依恋等积极情绪会促使用户在遇到问题时表现出更高的容

忍度与耐心，愿意克服困难并持续使用服务^[50-51]，从而降低了间歇性中辍的可能性^[52]。已有研究表明，情感承诺能有效降低用户的间歇性中辍行为倾向^[53-54]。因此，本文提出以下假设：

H7：情感承诺负向影响用户的间歇性中辍行为

3 研究设计

研究模型共包括 8 个变量，每个变量通过 3~4

项指标进行测量。这些指标选取自国内外文献中验证过的成熟量表，以增强量表在内容层面的有效性。问卷设计过程中，运用了翻译与反向翻译技术确保其准确无误。研究邀请了 15 名活跃的生成式 AI 用户参与预测试，并根据他们的反馈意见，对部分指标进行了调整优化。最终的测量指标及其文献来源如表 3 所示。

表 3 变量及测量指标
Tab. 3 Variables and Measurement Items

变 量	指标	指标内容	来源
隐私担忧 (PC)	PC1	我担心使用生成式 AI 会泄露个人隐私	[55]
	PC2	我担心生成式 AI 会过度收集我的个人信息	
	PC3	我担忧生成式 AI 会不合理使用我的个人信息	
信息幻觉 (IH)	IH1	生成式 AI 提供的信息可能是其编造的	[56]
	IH2	生成式 AI 提供的信息不大可靠	
	IH3	生成式 AI 提供的信息可能是虚假的	
智能化 (PI)	PI1	我觉得生成式 AI 具备专业能力	[36]
	PI2	我觉得生成式 AI 知识渊博	
	PI3	我觉得生成式 AI 具有较高的智能性	
拟人化 (PA)	PA1	生成式 AI 像人一样进行交互	[41]
	PA2	生成式 AI 是有意识的	
	PA3	生成式 AI 像有生命一样	
个性化 (PP)	PP1	我已经调整了生成式 AI 的设置以满足我的需要(如要求 AI 从特定领域作答)	[57]
	PP2	我可以选择生成式 AI 的功能，以适应我的使用风格(如对话样式中选择创造力、平衡或精确)	
	PP3	生成式 AI 提供的信息符合我的需求	
认知失调 (CD)	CD1	有时我有点后悔使用生成式 AI(如对个人信息泄露的担忧)	[46]
	CD2	使用生成式 AI 有时让我不舒服(如虚假新闻信息、虚假学术信息、偏见与歧视等)	
	CD3	当使用生成式 AI 的时候，我有时感到矛盾(如，完成任务时借助生成式 AI 又害怕过度依赖)	
	CD4	使用生成式 AI 有时让我对自己的信念(如人生观、价值观、习惯等)产生质疑	
情感承诺 (AC)	AC1	我很高兴去使用生成式 AI	[49]
	AC2	我喜欢去使用生成式 AI	
	AC3	使用生成式 AI 对我很重要	
	AC4	使用生成式 AI 是必要的，且符合我的偏好	
间歇性中辍 (ID)	ID1	我可能在未来降低使用生成式 AI 的频率	[23]
	ID2	我曾经停止使用生成式 AI 一段时间后又重新使用	
	ID3	我会暂停使用生成式 AI，未来如果有需求我会再重新使用	

调查问卷设计与发放借助在线问卷平台完成，并通过社交媒体以二维码和链接的形式分发给目标用户群体。经过筛选排除了填写时间过短或未曾使

用过生成式 AI 的无效问卷后，最终收集到有效问卷 405 份。在性别分布上，男性为 44.9%，女性为 55.1%；78.5%的受访者年龄位于 40 岁以下；60%

的受访者拥有本科及以上学历背景。使用经验方面, 76.8%的受访者使用过国内诸如文心一言、通义千问等产品, 75.1%的受访者使用过 ChatGPT、New Being 等国外生成式 AI 产品。使用频率方面, 11.4%的受访者每日使用生成式 AI, 38%的受访者表示他们的使用频率为每隔几天使用 1 次。

4 数据分析

4.1 结构方程模型 (SEM)

4.1.1 信效度检验

采用 SPSS 25.0 和 AMOS 21.0 对测量模型进行分析, 以检验量表的信度与效度。根据表 4 的数据, 各变量的 Alpha 系数介于 0.778~0.910 之间, 均超过了 0.70 的阈值, 显示了较好的内部一致性。同时, 各 CR 值位于 0.776~0.912 之间, 且所有因子载荷均大于 0.70, AVE 也都高于 0.50 的阈值, 显示量表具有良好的收敛效度。此外, 如表 5 所示, 各变量的 AVE 平方根值均显著大于该变量与其他变量间的相关系数, 显示较好的区分效度。

4.1.2 假设检验

采用 AMOS 21.0 对结构模型进行分析, 计算模型中各变量间的路径系数并进行假设检验。图 2 显示各假设均得到支持。表 6 列出了模型拟合结果, 所有拟合指数的实际值均优于推荐阈值, 表明该模型整体具有良好的拟合优度。

4.2 模糊集定性比较分析 (fsQCA)

传统的回归分析认为自变量之间是相互独立的, 所以 SEM 在探究自变量对因变量的边际“净效应”

表 4 信效度检测					
Tab. 4 Reliability and Validity Test					
变 量	指标	标准负载	Alpha	CR	AVE
隐私担忧 (PC)	PC1	0.769	0.806	0.806	0.580
	PC2	0.777			
	PC3	0.739			
信息幻觉 (IH)	IH1	0.792	0.824	0.826	0.613
	IH2	0.789			
	IH3	0.767			
智能化 (PI)	PI1	0.855	0.825	0.831	0.621
	PI2	0.761			
	PI3	0.744			
拟人化 (PA)	PA1	0.879	0.831	0.840	0.637
	PA2	0.776			
	PA3	0.733			
个性化 (PP)	PP1	0.714	0.778	0.776	0.536
	PP2	0.746			
	PP3	0.736			
认知失调 (CD)	CD1	0.879	0.903	0.904	0.702
	CD2	0.823			
	CD3	0.823			
	CD4	0.825			
情感承诺 (AC)	AC1	0.890	0.893	0.895	0.681
	AC2	0.808			
	AC3	0.794			
	AC4	0.806			
间歇性中辍 (ID)	ID1	0.874	0.910	0.912	0.775
	ID2	0.902			
	ID3	0.864			

表 5 相关系数矩阵								
Tab. 5 Correlation Coefficients Matrix								
	PC	IH	PI	PA	PP	CD	AC	ID
PC	0.762							
IH	0.298	0.783						
PI	0.346	0.329	0.788					
PA	0.408	0.344	0.397	0.798				
PP	0.421	0.416	0.311	0.363	0.732			
CD	0.373	0.379	0.195	0.218	0.243	0.838		
AC	0.276	0.258	0.353	0.402	0.432	0.155	0.825	
ID	0.178	0.184	0.061	0.067	0.077	0.567	0.063	0.880

时虽有其价值, 却难以解决前因变量间潜在的相互关联性 & 非对称因果关系难题。而 fsQCA 方法借鉴

整体论与集合论原理, 有效应对了此类问题^[58]。该方法考察了由多个条件变量构成的组态对结果变量

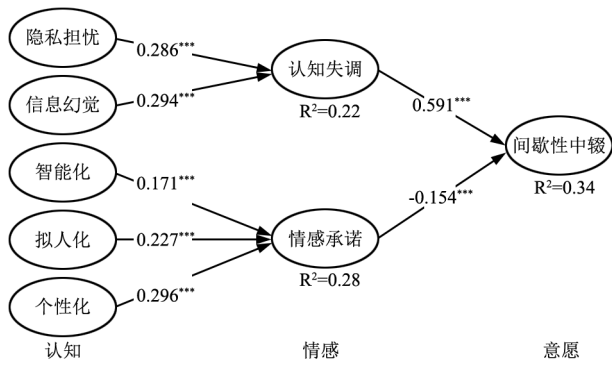


图 2 路径系数及显著性

Fig. 2 Path Coefficients and Significance

表 6 模型拟合指数

Tab. 6 Model Fit Indices

拟合指数	Chi²/df	GFI	AGFI	CFI	RMSEA
推荐值	<3	>0.90	>0.80	>0.90	<0.08
实际值	1.814	0.910	0.888	0.960	0.045

的影响，能够揭示变量之间的复杂因果关系^[59]。因此，本文采用 fsQCA 方法，研究导致生成式 AI 用户间歇性中辍的组态。

4.2.1 变量的选取与校准

本文选取 7 个前因变量，即隐私担忧、信息幻觉、智能化、拟人化、个性化、认知失调及情感承诺，将生成式 AI 用户的间歇性中辍设定为研究的结果变量。对所有变量值进行平均处理，随后遵循 Ragin C C^[60] 提出的完全不隶属阈值 5%、完全隶属阈值 95% 以及交叉点标准 50%，采用 Calibrate 函数对各个变量逐一实施数据校准。完成数据校准后，进一步进行了必要性条件分析。结果表明，各个前因变量单独作为结果变量的必要条件时，其一致性水平均未达到 0.9。这意味着没有任何一个单一的前因变量是引致生成式 AI 用户间歇性中辍行为的必要条件。因此，需要进一步采用组态分析来考察生成式 AI 用户间歇性中辍行为。

4.2.2 组态分析

在实施 fsQCA 分析时，首先构建了一张包含 2⁷ (128) 种可能条件组合的初始真值表，每行代表一个特定的条件搭配情景。结合样本的实际分布特征，本文将案例频数的可接受阈值设定为 7^[61]，一致性阈值设定为 0.8^[62]，PRI 阈值设为 0.7^[63]。表 7 展示了组态分析结果。其中，●表示核心条件存在，

⊗表示核心条件缺失，●表示辅助条件存在。

表 7 组态分析结果

Tab. 7 Configuration Analysis Results

条件变量	间歇性中辍	
	路径 1	路径 2
隐私担忧	●	⊗
信息幻觉	●	●
智能化	●	●
拟人化	●	●
个性化	●	●
认知失调	●	●
情感承诺	⊗	●
一致性	0.911	0.892
覆盖率	0.251	0.248
净覆盖率	0.063	0.060
总体一致性	0.891	
总体覆盖率	0.311	

路径 1 为“隐私担忧 * 信息幻觉 * 智能化 * 拟人化 * 个性化 * 认知失调 * ~情感承诺”，其中，隐私担忧和信息幻觉作为核心条件存在，智能化、拟人化、个性化和认知失调作为辅助条件存在，而情感承诺作为核心条件缺失。路径 1 揭示，在用户对生成式 AI 具有高度隐私担忧且信息幻觉及认知失调处于较高水平时，也就是使能因素处于较高值的情况下，即使包括智能化、拟人化和个性化在内的抑制因素水平较高，若情感承诺较低，则用户极可能表现出间歇性中辍的行为，从而突显出情感承诺在抑制间歇性中辍上的重要作用。

路径 2 为“~隐私担忧 * 信息幻觉 * 智能化 * 拟人化 * 个性化 * 认知失调 * 情感承诺”，其中，隐私担忧作为核心条件缺失，信息幻觉和认知失调作为核心条件存在，而智能化、拟人化、个性化和情感承诺作为辅助条件存在。路径 2 表明，在隐私担忧较轻的用户群体中，尽管使能因素(信息幻觉和认知失调)和抑制因素(智能化、拟人化、个性化和情感承诺)皆处于高位，但由于使能因素为核心条件，抑制因素为辅助条件，仍会诱发间歇性中辍。这类用户高度关注生成式 AI 的内容真实可靠性，由此而导致的认知失调将引发他们的间歇性中辍行为。

5 结果讨论

5.1 主要研究发现

1) 使能因素方面, 隐私担忧和信息幻觉显著影响生成式 AI 用户认知失调。fsQCA 也显示隐私担忧和信息幻觉是组态路径 1 的核心条件。一方面, 隐私担忧反映出个体对其个人信息可能被生成式 AI 不当收集、使用或泄露的深切忧虑。当用户对生成式 AI 技术的隐私保护机制持怀疑态度时, 这种疑虑会加剧他们对自身隐私权益可能受损的担忧^[29], 从而诱发一定程度的认知失调。另一方面, 信息幻觉反映了用户对 AI 生成内容的虚假性的感知, 深度伪造的信息可能使用户无法准确判断接收到的信息真伪, 造成认知失调状态。当用户因 AI 生成内容而产生错误信念或者做出不符合实际的决策时, 他们的认知系统便会在理想预期与实际体验之间产生不协调, 这将进一步加深用户的认知失调感受。

2) 抑制因素方面, 智能化、拟人化、个性化显著影响用户情感承诺。其中, 个性化的作用最大, 显示用户高度重视生成式 AI 的个性化内容和服务, 从而建立双方的关系连接, 这与已有文献结果相一致^[44]。fsQCA 分析显示智能化、拟人化、个性化是两条路径的共同辅助条件。首先, 智能化表现为系统的自主学习能力、决策优化以及问题解决的高效性, 这种能力增强了用户对其的信任感和依赖度, 进而促使用户形成强烈的情感连接。其次, 拟人化设计通过赋予生成式 AI 类似人类的特征和交互风格, 如自然语言处理、情绪感知及响应机制等, 使用户能够建立更深层次的人机连接。这种类人互动方式有助于增进用户的亲近感与认同感, 从而提高他们对系统的情感投入和承诺水平。最后, 个性化反映了生成式 AI 能够依据用户的行为习惯、兴趣偏好和需求进行动态适应与个性化服务提供。当用户感受到系统对其需求的理解和尊重时, 会更容易产生共情, 形成情感依恋和承诺。

3) 认知失调、情感承诺分别对用户间歇性中辍行为产生了显著的正向、负向影响, 且认知失调的作用强于情感承诺, 显示相对于抑制因素, 使能因素对间歇性中辍的作用更大。fsQCA 的路径 2 也显示认知失调是核心条件, 情感承诺是辅助条件, 这与 SEM 结果相一致。一方面, 内在的不和谐状态

促使用户在面对不确定性或不满意的服务体验时更倾向于选择暂时停止或减少使用该技术。当用户对生成式 AI 的心理预期与其实际使用体验存在较大差距时, 可能引发认知失调^[47], 用户间歇性中辍的可能性随之增加。另一方面, 当用户对生成式 AI 形成较强的情感连接, 如依赖感和认同感, 他们更可能持续使用并克服暂时的问题或困扰。然而, 情感承诺形成的过程通常较为缓慢且所需时间较长, 认知失调却可能导致用户迅速产生负面情绪和短期的适应性行为, 使认知失调对用户间歇性中辍的作用更为明显。

5.2 理论贡献

本文的理论贡献包括: ①已有研究主要关注生成式 AI 用户的积极行为(如采纳、持续使用等), 较少考察用户的消极行为(如间歇性中辍)。因此, 本文的研究结果丰富了关于生成式 AI 用户行为的研究。②基于 C-A-C, 本文发现认知因素包括隐私担忧、信息幻觉、智能化、拟人化、个性化, 影响用户的情感包括认知失调、情感承诺, 进而影响间歇性中辍行为。研究结果揭示了生成式 AI 用户间歇性中辍行为的形成机理。③本文从“使能”与“抑制”两个视角考察了生成式 AI 用户间歇性中辍行为。相对单一视角, 本文结果为理解生成式 AI 用户间歇性中辍提供更为完整的视角。④已有研究聚焦于负面情绪(如疲劳、焦虑和不满意等)对间歇性中辍的影响, 本文发现情感承诺是抑制生成式 AI 用户间歇性中辍的重要因素, 这一结果也充实了关于用户间歇性中辍的研究成果。

5.3 实践启示

本文的研究结果具有以下实践启示: ①缓解用户隐私担忧。生成式 AI 平台应重视数据安全和隐私保护措施, 通过设计更严格的数据加密机制、清晰的隐私政策说明以及可定制化的权限设置, 减轻用户在使用过程中的隐私顾虑。②减少信息幻觉。平台需要完善 AI 模型架构, 优化训练算法, 加强过程监督与反馈, 提高信息的真实可靠性, 减少 AI 幻觉问题, 以免引发用户认知失调。③降低用户的心理失调。平台需提高模型的透明度和可解释性, 帮助用户理解生成式 AI 的能力边界及其工作原理, 培养他们批判性地接收和评估 AI 生成信息的能力,

从而避免因误解或过度依赖 AI 输出而造成的心理压力。④增进用户的情感承诺。平台应提供个性化的用户体验,例如通过深度学习和用户画像技术提供更为贴近用户需求的服务,进而加深用户与平台之间的情感连接。

6 结 语

基于 C-A-C,本文考察了生成式 AI 用户间歇性中辍行为,采用 SEM 与 fsQCA 方法。研究发现用户间歇性中辍行为受到使能因素与抑制因素的双重影响。研究结果提升了对生成式 AI 用户间歇性中辍行为的理解,有助于生成式 AI 平台采取措施来抑制用户间歇性中辍行为,从而实现用户保持。

本文的局限包括:①生成式 AI 处于快速发展之中,其功能和应用场景不断丰富,并且从通用大模型发展到专业大模型。未来的研究可考察专业型 AI 用户中辍行为。②中辍行为是动态发展的过程,本文主要采用截面数据,未来的研究可采集纵向数据,跟踪用户中辍行为的发展。③本文主要考察了认知失调和情感承诺对间歇性中辍的影响,未来的研究可以考察其他变量如信任、用户体验等因素的作用。

参 考 文 献

- [1] 曹树金,曹茹桦.从 ChatGPT 看生成式 AI 对情报学研究与实践的影响[J].现代情报,2023,43(4):3-10.
- [2] 陈永伟.超越 ChatGPT:生成式 AI 的机遇、风险与挑战[J].山东大学学报(哲学社会科学版),2023,(3):127-143.
- [3] 张明新,叶银娇.传播新技术采纳的“间歇性中辍”现象研究:来自东西方社会的经验证据[J].新闻与传播研究,2014,21(6):78-98,127-128.
- [4] Margaret N Y M. Re-Examining the Innovation Post-Adoption Process: The Case of Twitter Discontinuance[J]. Computers in Human Behavior, 2020, 103: 48-56.
- [5] Menon D, Shilpa K. “Chatting with ChatGPT”: Analyzing the Factors Influencing Users’ Intention to Use the Open AI’s ChatGPT Using the UTAUT Model[J]. Heliyon, 2023, 9(11): e20962.
- [6] Li W Y. A Study on Factors Influencing Designers’ Behavioral Intention in Using ChatGPT-Generated Content for Assisted Design: Perceived Anxiety, Perceived Risk, and UTAUT[J]. International Journal of Human-Computer Interaction, 2024: 1-14.
- [7] Faruk L I D, Rohan R, Nimrutsirikun U, et al. University Students’ Acceptance and Usage of Generative AI (ChatGPT) from a Psycho-Technical Perspective [C] //Proceedings of the 13th International

- Conference on Advances in Information Technology, 2023: 1-8.
- [8] Ma X Y, Huo Y D. Are Users Willing to Embrace ChatGPT? Exploring the Factors on the Acceptance of Chatbots from the Perspective of AIDUA Framework [J]. Technology in Society, 2023, 75: 102362.
 - [9] 毛太田,汤淦,马家伟,等.人工智能生成内容(AIGC)用户采纳意愿影响因素识别研究——以 ChatGPT 为例[J].情报科学,2024,42(7):126-136.
 - [10] 张海,刘畅,王东波,等. ChatGPT 用户使用意愿影响因素研究[J].情报理论与实践,2023,46(4):15-22.
 - [11] Pham H C, Duong C D, Nguyen G K H. What Drives Tourists’ Continuance Intention to Use ChatGPT for Travel Services? A Stimulus-Organism-Response Perspective [J]. Journal of Retailing and Consumer Services, 2024, 78: 103758.
 - [12] Makkonen M, Salo M, Pirkkalainen H. The Effects of Job and User Characteristics on the Perceived Usefulness and Use Continuance Intention of Generative Artificial Intelligence Chatbots at Work [C] //CEUR Workshop Proceedings, 2023: 103-120.
 - [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] Rogers E M. Diffusion of Innovations [M]. New York: Free Press, 2005.
 - [15] 祝建华,何舟.互联网在中国的扩散现状与前景:2000 年京、穗、港比较研究[J].新闻大学,2002,(2):23-32.
 - [16] 甘春梅,林晶晶,肖晨.社交网络用户间歇性中辍行为的影响因素研究[J].情报理论与实践,2021,44(1):118-123.
 - [17] 张敏,孟蝶,张艳. S-O-R 分析框架下的强关系社交媒体用户中辍行为的形成机理——一项基于扎根理论的探索性研究[J].情报理论与实践,2019,42(7):80-85,112.
 - [18] Zhang S W, Zhao L, Lu Y B, et al. Do You Get Tired of Socializing? An Empirical Explanation of Discontinuous Usage Behaviour in Social Network Services [J]. Information & Management, 2016, 53(7): 904-914.
 - [19] 李君君,王金歌,曹园园.间歇性中辍行为特征的探索性研究——以微博数据为例[J].现代情报,2021,41(1):60-66.
 - [20] 杜诗瑶.基于扎根理论的短视频平台用户中辍行为的影响因素研究[J].东南传播,2021,(12):96-100.
 - [21] 甘春梅,明昕宇. SOR 视角下短视频用户间歇性中辍意愿的实证研究[J].情报科学,2023,41(6):153-160.
 - [22] 王文韬,宋天晓,唐思捷,等.短视频用户间歇性中辍行为反复机理及优化研究[J].现代情报,2023,43(12):51-62.
 - [23] 沈校亮,厉洋军.智能健康硬件用户间歇性中止行为影响因素研究[J].管理科学,2017,30(1):31-42.
 - [24] Dai B, Ali A, Wang H W. Exploring Information Avoidance Intention of Social Media Users: A Cognition-Affect-Conation Perspective [J]. Internet Research, 2020, 30(5): 1455-1478.
 - [25] Huang Y M, Lou S J, Huang T C, et al. Middle-Aged Adults’

- Attitudes Toward Health App Usage: A Comparison with the Cognitive-Affective-Conative Model [J]. Universal Access in the Information Society, 2019, 18 (4): 927-938.
- [26] Martínez-Navarro J, Bigné E, Guixeres J, et al. The Influence of Virtual Reality in E-Commerce [J]. Journal of Business Research, 2019, 100: 475-482.
- [27] 彭丽徽, 李贺, 张艳丰, 等. 用户隐私安全对移动社交媒体倦怠行为的影响因素研究——基于隐私计算理论的 CAC 研究范式 [J]. 情报科学, 2018, 36 (9): 96-102.
- [28] Guo Y Y, Wang X, Wang C Y. Impact of Privacy Policy Content on Perceived Effectiveness of Privacy Policy: The Role of Vulnerability, Benevolence and Privacy Concern [J]. Journal of Enterprise Information Management, 2022, 35 (3): 774-795.
- [29] 陈昊, 李文立, 柯育龙. 社交媒体持续使用研究: 以情感响应为中介 [J]. 管理评论, 2016, 28 (9): 61-71.
- [30] Hatem R, Simmons B, Thornton J E. A Call to Address ChatGPT “Hallucinations” and How Healthcare Professionals Can Mitigate Their Risks [J]. Cureus, 2023, 15 (9): e44720.
- [31] 王松, 王瑜, 李芳. 匹配视角下社会化商务用户消极使用行为形成机理研究——基于认知失调的中介 [J]. 软科学, 2020, 34 (10): 133-139.
- [32] Wang R, He Y, Xu J, et al. Fake News or Bad News? Toward an Emotion-Driven Cognitive Dissonance Model of Misinformation Diffusion [J]. Asian Journal of Communication, 2020, 30 (5): 317-342.
- [33] 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.
- [34] Moussawi S, Koufaris M. Perceived Intelligence and Perceived Anthropomorphism of Personal Intelligent Agents: Scale Development and Validation [C] //Proceedings of the 52nd Hawaii International Conference on System Sciences, 2019: 115-124.
- [35] Song X, Xu B, Zhao Z Z. Can People Experience Romantic Love for Artificial Intelligence? An Empirical Study of Intelligent Assistants [J]. Information & Management, 2022, 59 (2): 103595.
- [36] Rafiq F, Dogra N, Adil M, et al. Examining Consumer's Intention to Adopt ChatGPT-Chatbots in Tourism Using Partial Least Squares Structural Equation Modeling Method [J]. Mathematics, 2022, 10 (13): 2190.
- [37] Moussawi S, Koufaris M, Benbunan-Fich R. How Perceptions of Intelligence and Anthropomorphism Affect Adoption of Personal Intelligent Agents [J]. Electronic Markets, 2021, 31 (2): 343-364.
- [38] Lin H X, Chi O H, Gursoy D. Antecedents of Customers' Acceptance of Artificially Intelligent Robotic Device Use in Hospitality Services [J]. Journal of Hospitality Marketing & Management, 2019, 29 (5): 530-549.
- [39] Pelau C, Dabija D C, Ene I. What Makes an ChatGPT Device Human-Like? The Role of Interaction Quality, Empathy and Perceived Psychological Anthropomorphic Characteristics in the Acceptance of Artificial Intelligence in the Service Industry [J]. Computers in Human Behavior, 2021, 122: 106855.
- [40] Singh R. “Hey Alexa-Order Groceries for Me” – the Effect of Consumer-Vai Emotional Attachment on Satisfaction and Repurchase Intention [J]. European Journal of Marketing, 2022, 56 (6): 1684-1720.
- [41] Zhang A D, Patrick Rau P L. Tools or Peers? Impacts of Anthropomorphism Level and Social Role on Emotional Attachment and Disclosure Tendency Towards Intelligent Agents [J]. Computers in Human Behavior, 2023, 138: 107415.
- [42] Li C. When Does Web-Based Personalization Really Work? The Distinction between Actual Personalization and Perceived Personalization [J]. Computers in Human Behavior, 2016, 54: 25-33.
- [43] Choi N. Information Systems Attachment: An Empirical Exploration of Its Antecedents and Its Impact on Community Participation Intention [J]. Journal of the American Society for Information Science and Technology, 2013, 64 (11): 2354-2365.
- [44] Sheng H, Xiao H. Examining Users' Continuous Use Intention of ChatGPT-Enabled Online Education Applications [C] //2022 International Conference on Computer Engineering and Artificial Intelligence (ICCEAI). IEEE, 2022, 642-645.
- [45] Festinger L. A Theory of Cognitive Dissonance [M]. New York: Stanford University Press, 1957.
- [46] Vaghefi I, Qahri-Saremi H, Turel O. Dealing with Social Networking Site Addiction: A Cognitive-Affective Model of Discontinuance Decisions [J]. Internet Research, 2020, 30 (5): 1427-1453.
- [47] 陈婷, 李霞, 段尧清. 短视频社交媒体用户不持续使用意向研究——整合认知失调与自我效能双重视角 [J]. 情报杂志, 2022, 41 (10): 199-207.
- [48] Allen N J, Meyer J P. The Measurement and Antecedents of Affective, Continuance and Normative Commitment to the Organization [J]. Journal of Occupational Psychology, 1990, 63 (1): 1-18.
- [49] 付少雄, 袁秋, 邓胜利, 等. 承诺理论视角下辟谣短视频分享意愿的驱动因素 [J]. 图书馆论坛, 2024, 44 (7): 128-138.
- [50] Cao Y Y, Long Q Q, Hu B L, et al. Exploring Elderly Users' Msn's Intermittent Discontinuance: A Dual-Mechanism Model [J]. Telematics and Informatics, 2021, 62: 101629.
- [51] Feng G C, Su X L, He Y R. A Meta-Analytical Review of the Determinants of Social Media Discontinuance Intentions [J]. Mass Communication and Society, 2024, 27 (3): 525-550.
- [52] 乐承毅, 陈征. PPM 理论框架下短视频用户非持续使用意向成因研究 [J]. 现代情报, 2022, 42 (6): 80-93.

(下转第 64 页)

- Counter Confirmation Bias [J]. *Frontiers in Big Data*, 2019, 2: 11.
- [35] Zhang K Z K, Cheung C M K, Lee M K O. Examining the Moderating Effect of Inconsistent Reviews and its Gender Differences on Consumers' Online Shopping Decision [J]. *International Journal of Information Management*, 2014, 34 (2): 89-98.
- [36] Jiang S, Beaudoin C E. Health Literacy and the Internet: An Exploratory Study on the 2013 HINTS Survey [J]. *Computers in Human Behavior*, 2016, 58: 240-248.
- [37] Tracey J, Arroll B, Barham P, et al. The Validity of General Practitioners' Self Assessment of Knowledge: Cross Sectional Study [J]. *BMJ*, 1997, 315 (7120): 1426-1428.
- [38] Angst C M, Agarwal R. Adoption of Electronic Health Records in the Presence of Privacy Concerns: The Elaboration Likelihood Model and Individual Persuasion [J]. *MIS Quarterly*, 2009, (2): 339-370.
- [39] Petty R E, Cacioppo J T, Goldman R. Personal Involvement as A Determinant of Argument-Based Persuasion [J]. *Journal of Personality and Social Psychology*, 1981, 41 (5): 847-855.
- [40] Kruglanski A W, Thompson E P. Persuasion by a Single Route: A View from the Unimodel [J]. *Psychological Inquiry*, 1999, 10 (2): 83-109.
- [41] 李松涛, 张卫东, 陈希鹏. 辟谣信息对社交媒体用户错误信息持续影响效应(CIE)的影响——基于知识修正理论 [J]. *现代情报*, 2024, 44 (5): 123-139.
- [42] Yang S B, Lee H, Lee K, et al. The Application of Aristotle's Rhetorical Theory to the Sharing Economy: An Empirical Study of Airbnb [J]. *Journal of Travel & Tourism Marketing*, 2018, 35 (7): 938-957.
- [43] Chou C H, Wang Y S, Tang T I. Exploring the Determinants of Knowledge Adoption in Virtual Communities: A Social Influence Perspective [J]. *International Journal of Information Management*, 2015, 35 (3): 364-376.
- [44] Wu P C S, Wang Y C. The Influences of Electronic Word-of-Mouth Message Appeal and Message Source Credibility on Brand Attitude [J]. *Asia Pacific Journal of Marketing and Logistics*, 2011, 23 (4): 448-472.
- [45] Langfred C W. The Downside of Self-Management: A Longitudinal Study of the Effects of Conflict on Trust, Autonomy, and Task Interdependence in Self-Managing Teams [J]. *Academy of Management Journal*, 2007, 50 (4): 885-900.
- [46] Preacher K J, Hayes A F. Asymptotic and Resampling Strategies for Assessing and Comparing Indirect Effects in Multiple Mediator Models [J]. *Behavior Research Methods*, 2008, 40 (3): 879-891.
- [47] 温忠麟, 侯杰泰, 张雷. 调节效应与中介效应的比较和应用 [J]. *心理学报*, 2005, 37 (2): 268-274.
- [48] 杨颖兮, 喻国明. 传播中的非理性要素: 一项理解未来传播的重要命题 [J]. *探索与争鸣*, 2021, (5): 131-138, 179.
- [49] Schwarz N, Sanna I J, Skurnik I, et al. Metacognitive Experiences and the Intricacies of Setting People Straight: Implications for Debiasing and Public Information Campaigns [J]. *Advances in Experimental Social Psychology*, 2007, 39: 127-161.
- [50] 陈忆金, 潘沛. 健康类短视频信息有用性感知的影响因素研究 [J]. *现代情报*, 2021, 41 (11): 43-56.
- (责任编辑: 杨丰侨)

.....

(上接第 50 页)

- [53] Yao X S, Ma S F, Wu Y, et al. So Said, So Done? The Role of Commitment in Activity-Based Check-in Discontinuance on Apps [J]. *Technological Forecasting and Social Change*, 2023, 194: 122675.
- [54] Wang J K, Zheng B W, Liu H F, et al. A Two-Factor Theoretical Model of Social Media Discontinuance: Role of Regret, Inertia, and Their Antecedents [J]. *Information Technology & People*, 2021, 34 (1): 1-24.
- [55] Rosenthal S, Wasenden O C, Gronnevet G A, et al. A Tripartite Model of Trust in Facebook: Acceptance of Information Personalization, Privacy Concern, and Privacy Literacy [J]. *Media Psychology*, 2020, 23 (6): 840-864.
- [56] Zha X J, Yang H J, Yan Y L, et al. Exploring the Effect of Social Media Information Quality, Source Credibility and Reputation on Informational Fit-to-Task: Moderating Role of Focused Immersion [J]. *Computers in Human Behavior*, 2018, 79: 227-237.
- [57] Cao X F, Gong M C, Yu L L, et al. Exploring the Mechanism of Social Media Addiction: An Empirical Study from WeChat Users [J]. *Internet Research*, 2020, 30 (4): 1305-1328.
- [58] 徐广平, 张金山, 杜运周. 环境与组织因素组态效应对公司创业的影响——一项模糊集的定性比较分析 [J]. *外国经济与管理*, 2020, 42 (1): 3-16.
- [59] 杜运周, 贾良定. 组态视角与定性比较分析(QCA): 管理学研究的一条新道路 [J]. *管理世界*, 2017, (6): 155-167.
- [60] Ragin C C. *Redesigning Social Inquiry: Fuzzy Sets and Beyond* [M]. Chicago: University of Chicago Press, 2008.
- [61] Pappas I O, Woodside A G. Fuzzy-Set Qualitative Comparative Analysis(fsQCA): Guidelines for Research Practice in Information Systems and Marketing [J]. *International Journal of Information Management*, 2021, 58: 102310.
- [62] Fiss P C. Building Better Causal Theories: A Fuzzy Set Approach to Typologies in Organization Research [J]. *Academy of Management Journal*, 2011, 54 (2): 393-420.
- [63] Greckhamer T, Furnari S, Fiss P C, et al. Studying Configurations with Qualitative Comparative Analysis: Best Practices in Strategy and Organization Research [J]. *Strategic Organization*, 2018, 16 (4): 482-495.
- (责任编辑: 杨丰侨)