

基于政务微博数据的人工智能生成应急信息感知质量研究^{*}

李一啸¹ 张璐¹ 周莎莎² 杨水清²

(1.浙江财经大学 信息技术与人工智能学院, 浙江 杭州 310018;

2.浙江财经大学 管理学院, 浙江 杭州 310018)

摘 要 应急信息的即时生产与传播对于突发事件的应急响应至关重要。在生成式人工智能广泛应用的背景下, 利用生成式人工智能的技术和工具能够提升应急信息的生产效率。本文爬取某次暴雨灾害期间的政务微博数据, 根据情境危机沟通理论将微博文本划分为三类, 通过隐含狄利克雷分布 (latent Dirichlet allocation, LDA) 模型提取各类信息的主题和词项, 并用于大语言模型的提示词优化, 从而生成相应的应急响应信息。在此基础上, 研究公众对于人类及人工智能生成信息的感知质量差异, 以及人口统计学特征对差异的影响。结果发现, 公众对于人类和人工智能生成的信息在感知质量上没有明显差异, 学历、年龄等因素却显著影响着不同类型信息的感知质量。本文探讨了人工智能生成技术和工具在应急信息生成过程中的应用可行性, 为政府、媒体机构等应急信息传播主体提供相关建议。

关键词 人工智能生成内容, 应急信息, LDA, 感知质量, ChatGPT

中图分类号 C939

1 引言

各类突发事件给人类生产生活带来严重的负面影响, 应急信息的发布和传播具有重要意义。以互联网为依托, 灵活高效的应急响应协作网络在社交媒体平台上涌现, 成为政府部门应对重大灾情时进行权威发布、服务救助以及舆论引导的重要渠道之一, 社交媒体允许用户进行内容创建与分享的功能特性也在灾害和危机信息管理中发挥了前所未有的作用^[1]。社交媒体的广泛应用使得应急信息传播的公众参与水平进一步提升, 但研究表明政府应急信息发布不及时、公众信息需求得不到满足等问题依旧突出^[2]。通过 2012~2021 年国际应急信息管理研究的词频统计发现^[3], 信息不对称、风险信息感知与沟通依然是社交媒体时代下应急信息管理领域的热点研究主题。危机信息学、信息系统等领域的研究者对此高度关注, 持续探讨如何应用不断迭代更新的信息技术解决这些难题。

人工智能的发展从学术研究到大规模应用, 一路突飞猛进。ChatGPT 的爆红进一步说明一类被称为大语言模型 (large language model, LLM) 的生成式人工智能应用取得了实质性进展^[4, 5]。近期的一些研究表明, 大语言模型 (以下简称大模型) 已经具备了应用于危机信息管理的潜在价值: 基于大模型构建的智能问答系统能够在紧急情形下成为媲美人类专家的智能顾问^[6]; 其在文本撰写和上下文感知提取方

^{*} 基金项目: 国家自然科学基金面上项目 (72371223)、教育部人文社会科学青年基金项目 (23YJC870015)。

通信作者: 张璐, 浙江财经大学信息技术与人工智能学院, 硕士研究生, E-mail: luz_1110@163.com。

面的能力使得信息处理过程能够高效进行,在极短时间内生成准确而全面的突发事件报告^[7]。人工智能技术生成内容的质量分析和评估吸引了众多学者的讨论。一方面,研究表明在诗歌创作^[8]、产品文案撰写^[9]等方面人工智能生成内容的质量已经达到了实际应用的要求;另一方面,也有研究发现由人工智能生成的公共卫生信息存在输出质量并不一致等问题,最终的内容质量仍然依赖于人类的干预和判断^[4]。在突发事件中,公众对应急信息的感知质量直接影响其对应急信息的采纳和应用程度,进而影响应急信息传播发挥的实际效果^[10]。

现有关于人工智能生成内容的用户感知质量研究为人工智能生成应急信息的应用奠定了基础,但仍存在以下不足:①尽管在非突发事件下的若干应用场景下人工智能生成内容的用户感知质量已经达到较为满意的水平^[8,9],但在公共卫生信息生成方面仍存在一定缺陷^[4],人工智能生成应急信息的质量是否达到有效辅助甚至替代人类的水平,仍需进一步检验。②尽管生成式人工智能技术已经被应用到危机信息管理中^[6,7],但危机情境作用、对用户信息需求的精准匹配依然容易被忽视,应急信息在保证内容准确的基础上更需要适应危机情境的动态演变^[11],生成式人工智能技术如何融合危机情境特征仍需深入探讨。③人工智能生成内容的用户感知质量测量是一个难点^[4,8,9],信息系统等领域中信息质量等相关变量的测量模型对此有借鉴价值。综上,有必要探究人工智能生成应急信息的感知质量及其与人类生成信息的差异,通过深入理解公众如何认知和评估人工智能生成的应急信息,为相关理论框架的完善提供实证支持。同时,弥补现有研究在危机情境下人工智能生成信息使用行为方面的不足,为设计人机协同的应急信息生成和优化策略提供依据,提升生成式人工智能技术在危机沟通中的应用效果。

基于此,本文选取某次极端暴雨案例,爬取若干重要政务微博账号发布的相关应急信息,基于情境危机沟通理论对原始信息进行主题提取,结合各类主题对应的关键词项优化大模型提示词,得到有针对性的人工智能生成的应急信息。本文结合情境危机沟通理论考虑应急信息不同类型,详细探究不同类型信息对用户感知质量的影响。参考信息系统等领域相关文献^[12-22]设计用户感知质量的测量模型,研究公众对人工智能(vs.人类)产生的不同类型应急信息的感知质量及差异,探讨人工智能生成应急信息的应用可行性。此外,近期的研究表明年龄、性别、社会地位等因素会影响人们对人工智能生成文本的判断^[23]。本文进一步结合人口统计学特征对公众的人工智能生成应急信息感知质量进行研究,剖析公众的人口统计学特征对上述差异的影响。综上,本文将回答以下两个研究问题:①公众对人工智能生成应急信息的感知质量如何以及相较于人类生成的信息是否存在差异?②此类差异在不同人群中是否会有所不同?

2 文献综述

2.1 危机信息管理与社交媒体

公共危机事件是任何社会发展阶段都要面对的问题,与每个社会成员息息相关。世界各国,每年因为自然灾害、突发事件、社会安全等危机事件造成的严重伤亡和经济损失数量巨大^[24]。随着社交媒体的广泛渗透以及移动互联网的广泛应用,社交媒体成为应对突发危机事件的重要信息平台之一,越来越多的学者开始研究危机信息管理与社交媒体响应的协同作用。社交媒体信息传播的各阶段彼此重叠^[25],对时效性的要求也越发严格。很早便有学者提出危机状态下政府的执政能力越来越表现为对信息的掌控及运用能力,有效的信息获取能力、成功的信息管理模式对于应对公共危机事件起着决定性作用^[26]。Wang 等^[27]分析了 2012 年北京暴雨期间不同时间段不同主题的微博占比变化,证明了社交媒体上

及时的应急信息可以加快应急事件的响应速度。Li 等^[28]对彝良地震灾害期间微博转发行为形成的动态网络进行了研究,通过机器学习方法将微博信息分类,揭示了灾害背景下不同类型信息的网络传播特征。Ogie 等^[29]系统梳理并总结了灾后恢复阶段社交媒体信息对于捐赠、灾后重建、信息支持、心理健康和情感支持等多个方面的覆盖程度与实际效果,并对社交媒体在灾后恢复中发挥更大作用的前景进行了展望。

社交媒体在危机信息管理中具有重要作用,为应急信息传播提供了新的渠道。一方面,它在提升信息传播效率与政民互动水平方面具有显著优势;另一方面,信息质量和公众信息需求得不到及时满足等问题依旧突出。现有研究较多关注危机状态下如何发挥社交媒体的工具性作用、政府对应急信息传播的治理等议题^[27, 29, 30],但仍需进一步探索如何提升应急信息的质量与发布效率。

2.2 生成式人工智能在危机信息管理中的应用

人工智能生成内容 (artificial intelligence generated content, AIGC) 通常被定义为通过生成式人工智能技术自动生成数字内容的新型内容生产方式^[31]。李白杨等^[32]总结了 AIGC 的技术特征,包括但不限于数据巨量化、内容创造力、跨模态融合、认知交互力等。在危机信息管理等领域, AIGC 技术的应用展现出强大的赋能潜力。在危机信息传播中,相关人员能够利用 AIGC 技术加速应急信息的生成与处理^[4]。在应急响应中,智能对话系统能够帮助不同类型的用户获取应急知识、获得决策指导等^[6],并能通过将 AIGC 整合至应急协作系统,为受危机影响的公众提供人工智能驱动的实时指令和信息^[33]。在应急情报服务中, AIGC 能够提升各个工作流程的效率^[34]。在风险分析与评估中,基于 AIGC 技术生成灾害报告能够在降低报告撰写时间开销的同时,保持较高的准确率^[7, 35]。在应急决策中,相关管理人员可以借助 AIGC 等技术实时解读来自紧急电话和社交媒体的数据,并据此采取行动^[36]。沙勇忠等^[37]对公共危机信息管理研究进行了文献回顾,并展望了生成式人工智能、AIGC、大模型等新技术涌现对该领域产生的影响,呼吁探索人与技术的协同。

人工智能生成应急信息在危机信息传播、应急响应、应急情报服务等方面具有重要应用价值,能够提高相关管理人员在突发事件下信息处理与发布的效率,解决信息发布不够及时、互动水平较低^[2]等问题,减轻相关人员的工作负担。现有的相关研究依靠机器学习、自然语言处理、神经重排序模型等理论方法,聚焦于业务流程的优化^[6, 7]、不同大模型的训练与改进^[33, 35, 36]等议题,旨在提高应急响应的效率,然而,这些研究侧重于技术层面的优化,对危机情境的动态演变关注不足,也较少从信息受众视角开展分析。在实际应急场景中,社交媒体上应急信息的接受者主要是普通民众,突发事件的发展态势处于持续变化中,人工智能生成的应急信息是否符合公众的心理预期有待探讨。

2.3 人工智能生成内容的感知质量

感知信息质量是指个人对其从发送者处接收到的信息的价值程度的感知^[38]。在数字新闻领域,人工智能生成的内容质量已达到一定水准,被用来标准化生成新闻片段^[39]。在应急信息传播的场景下,信息发布的权威性常常被视为影响感知质量的关键因素^[40],尤其当信息来源过于多元时,公众往往倾向于选择政府或专业机构发布的信息。在内容维度上,苏君华和杜念^[41]梳理了准确性、及时性、完整性、一致性等应急信息质量评价指标,Seppänen 和 Virrantaus^[42]认为及时、准确是应急信息质量的重要体现。Zhou 等^[43]从信息评估的角度,使用大模型生成了新冠疫情背景下的谣言,发现人工智能生成的错误信息倾向于满足证据可信、来源透明等质量方面的标准,给人们辨别信息带来了巨大的困扰。也有研究表明,生成式人工智能技术需要适应不同的输入以及背景,仅仅依靠大模型生成的内容往往缺乏必要的灵活性^[44],未能更好地捕捉和调节不同情境下公众的需求与反应。

用户感知质量是 AIGC 使用行为研究中的关键问题之一,直接影响公众对生成式人工智能技术应用的接受程度和信任程度。公众对信息质量的评估受到来源可信度^[10]、信息透明度^[43]、个人情感^[9]等多重因素的影响。Dwivedi 等 70 多位来自信息系统等多个领域的学者共同探讨了生成式对话 AI(如 ChatGPT)在研究、实践和政策方面的影响^[45],文中多次提到了对人工智能生成内容质量的担忧。尽管现有研究结合不同领域、不同应用场景开展了人工智能和人类生成信息的差异分析,然而结合社交媒体、应急信息传播两方面背景的研究相对较少,并且缺乏详细的测量指标来刻画用户对相关信息的感知质量。在应急管理这一特殊的应用场景下,人工智能生成应急信息的质量影响应急实践活动的最终效果,是实现“负责任的人工智能”应用的关键因素。对此,本文针对社交媒体应急信息传播这一应用场景,设计基于大模型的应急信息生成框架,研究公众对于人工智能生成信息的感知质量,通过对人、机信息质量的对比分析,补充用户信息使用行为在该部分的空白。其中,对用户感知质量的测量参考了信息系统等领域关于信息质量研究的相关文献。

2.4 情境危机沟通理论

情境危机沟通理论(situational crisis communication theory)是美国学者 Coombs^[46]提出的一个理论框架,解释了组织在应对不同危机时如何选择相应的应对策略,从而达到挽救组织声誉的效果^[47]。其主要研究两个方面的内容,一是危机情境,二是反应策略。情境危机沟通理论给出了较为详细的使用危机应对策略的建议以及匹配说明。该理论指出在危机沟通的开始便关注组织的声誉是不负责任的,必须先通过沟通(使用指导信息和调整信息)来解决受害者的生理和心理问题,这是一种道德责任。在此基础上,可进一步参考该理论指导应对策略的设计和应用^[48],这一思路已经在各级政府的社交媒体中有了不同程度的体现^[49]。

情境危机沟通理论强调组织应根据危机责任归属、严重性及公众情绪等情境要素,动态选择沟通策略以修复信任,而用户感知质量则从信息接收者视角评估信息的有效性、可信度与满意度。在实际情况中,组织可能采取各种策略应对危机,这些策略决定了信息传播的内容和方式,影响受众对危机信息质量的感知,尤其是对信息类型差异化的感知。传统的应急信息发布机制往往不仅缺乏对信息的动态管理,还对相关人员提出了较高的要求^[50]。基于情境危机沟通理论,本文将应急信息划分为指导信息、调整信息与支持信息三大类,进一步对微博数据进行主题提取并按照不同应急信息类型生成具备不同功能的信息,灵活适应危机的动态演化,优化信息生成策略,提高应急信息质量。

2.4.1 指导信息

指导信息最主要的作用是告诉公众他们应该采取什么行动来保护自己免受危机的伤害^[51]。对于大多数处于危机中的组织来说,指导信息都应先于修复组织声誉的策略。已有研究证实,指导信息的作用对声誉的影响起决定性作用^[52]。该类信息包括对危机内容、成因、时间、地点的说明以及为使危害最小化而采取的预防或纠正措施^[47]。该类信息一般具有权威性,包含基本事实,准确及时是其重要标准^[42],信息的接收者通常会依据信息是否能提供切实可行的指导对其进行评估。在自然灾害中,经常使用指导信息策略,这些信息可能采取实时灾难更新、救援报告、旅行建议等形式^[53]。

2.4.2 调整信息

调整信息有助于公众面对危机时进行心理层面的调整与应对,通过传播调整信息,组织表达对受危机影响的人们的关注^[51]。调整信息的一个重要组成部分是对受害者表达关心,这不仅有助于减少危机的负面影响,也是利益相关者所期望的^[48]。实证研究发现,调整信息策略往往包含同情、希望等情绪表达,

使用调整信息有利于维持灾后恢复的希望^[47]。调整信息需要结合受众背景与需求，否则可能影响信息的实际效用。研究表明，在危机沟通策略中，该类信息通常与指导信息结合使用^[54]；但又与指导信息不同，其作用更多体现在情绪关怀。调整信息的语言风格、情绪表达都会影响信息的接受度和可信度。Page^[55]构建的危机期间指导信息和调整信息质量评估量表表明，情感支持的提供是衡量调整信息质量的重要维度之一。

2.4.3 支持信息

支持策略在自然灾害方面可以表现为对相关负责人付出的肯定以及对受害者的同情，这一策略有利于在灾难恢复阶段鼓舞士气，为受害者带来信心^[47]。Coombs^[48]指出，与利益相关者有积极关系的管理者可以利用善意支持来帮助维护组织声誉。赞扬利益相关者在危机期间的努力，可以作为改善与他们关系的一种手段。政府组织也会策略性地让媒体和社会成员参与进来，通过强调组织的正面形象来积极恢复集体认同，这需要该类信息体现出赞扬与信任之意。相较于表达准确、以预防提示为主的指导信息，表达积极、能够引发正能量的信息往往被用于支持策略中^[54]，从而给用户带来更高的感知价值。在此次爬取的应对灾害的微博信息中也发现了支持策略的运用，因此在本文中将使用该类策略的微博文本称为支持信息。

3 研究方法与设计

3.1 数据采集与 AI 生成应急信息

研究的第一项内容是基于政务微博数据的应急响应信息生成。大致的技术路线为，根据情境危机沟通理论，对爬取的政务微博文本进行分类，对每类信息采用 LDA 模型提取主题，根据提取到的信息主题及相关词项对大模型（本文采用 ChatGPT）的提示词进行优化，最后通过与 ChatGPT 的对话输出由人工智能创造的应急响应信息。该过程如图 1 所示。

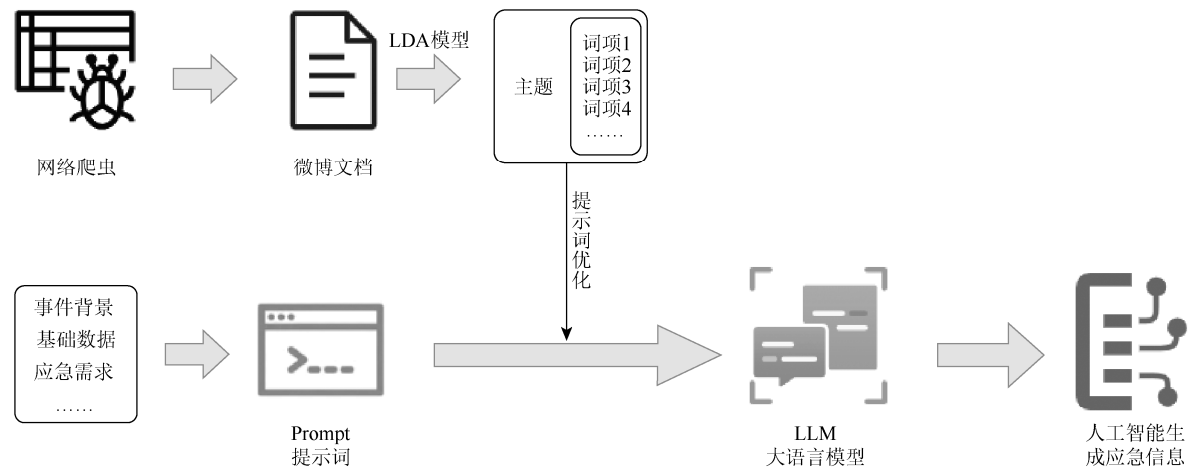


图 1 应急信息生成过程

3.1.1 微博数据收集

根据 7 月 31 日至 8 月 10 日期间政务微博排行榜以及发博量，挑选出如平安北京、北京发布、河北

天气、河北消防等八个暴雨洪灾中发博较为活跃的市级及以上政务微博账号，爬取其与雨情相关的原创微博内容。由于本次研究只涉及文字内容生成，去除部分纯图像、视频微博后，最终得到 1449 条微博作为本次研究信息集。

3.1.2 信息分类及主题提取

为了更好地呈现人类生成的应急响应信息，根据上文提到的响应策略中三类不同的信息类型，由两位研究人员独立地对爬取的微博文本进行分类，整合意见后完成信息分类，进而提取不同类型下信息的主题。本文选择 LDA 主题建模来分析不同类型信息的主题，该方法基于贝叶斯概率模型，是一种典型的词袋模型，利用词频的共现频率进行词组聚类^[56]。由于 LDA 不能确定最佳主题数目，需要先确定主题个数后对模型进行训练。一般用来评价 LDA 主题模型的指标有困惑度和主题一致性^[57, 58]。其中，困惑度被用来衡量一个模型在预测下一个单词时的不确定性。如果一个模型能够准确地预测下一个单词，那么它的困惑度就会比较低，当困惑度下降时，代表模型的预测越来越准确。但是，当主题数很多或存在其他原因时，生成的模型往往会过拟合，因此本文同时采用主题一致性与困惑度作为评估指标。由于 LDA 主题抽取的效果与主题数目的确定有直接关系^[58]，考虑到实际情况，在计算时先设置主题数 K 的范围为 0~20。接下来观察该范围内两种计算方法下主题个数最优值，最终根据得分结果确定每类信息下的最佳主题个数为 6。利用 Python 可视化工具 LDAvis 包对主题下的特征词进行可视化分析，根据高概率词项总结出的主题以及词项分布见表 1~表 3。

表 1 指导信息主题词项分布

主题	主题标识	与主题相关的 15 个高频词项
1	地域气象预警	地区、保定、张家口、天气、北部、承德、短时、西部、河北省、预警、石家庄、局地、阵雨、邢台、夜间
2	降水预警	毫米、降雨、小时、暴雨、平均、天气、雨强、气象台、北京、地区、雷阵雨、发布、局地、预计、实况
3	灾害预警与防护提示	预警、暴雨、风险、北京、防汛、气象、发布、地质灾害、红色、提示、发生、灾害、洪水、城市、响应
4	公共防灾措施	北京、影响、房山、台风、城区、小时、昌平、全市、丰台、工作、东北、水库、杜苏芮、门头沟、通信
5	地域交通预警与管控措施	河北、高速、通行、路段、卡努、台风、服务台、城市、车辆、交警、措施、暴雨、断交、北京、风险
6	城区交通管理与道路通行	断交、积水、路段、大街、地道、通行、绕行、石家庄、台风、城市、双向、卡努、服务台、河北、南三环

表 2 调整信息主题词项分布

主题	主题标识	与主题相关的 15 个高频词项
1	官兵救援与群众安抚	被困、救援、河北、暴雨、北京、群众、消防员、人员、老人、洪水、积水、转移、交警、城市、台风
2	中央指挥与救援关注	防汛、工作、救灾、群众、北京、救援、恢复、转移、物资、重建、人员、房山区、习近平、检查、灾后
3	通信恢复与物资运送	北京、通信、物资、应急、救援、无人机、工作、恢复、道路、运送、门头沟区、车辆、保障、门头沟、房山区
4	应急救助与政策调整	救援、灾情、缴存、公积金、人员、住房、消防、影响、措施、职工、北京、提取、相关、气象、防汛

续表

主题	主题标识	与主题相关的 15 个高频词项
5	群众转移与救援调整	北京、救援、平安、涿州、暴雨、房山、消防员、积水、人员、门头沟、消防、河北、风雨、被困、民警
6	应急救援与水患恢复	断交、积水、路段、大街、地道、通行、绕行、石家庄、台风、城市、双向、卡努、服务台、河北、南三环

表 3 支持信息主题词项分布

主题	主题标识	与主题相关的 15 个高频词项
1	交通人员救援与坚守	群众、北京、刘捷、暴雨、平安、乘客、中门、转移、工作、次列车、险情、救援、致敬、孩子、滞留
2	群众协作与救援互助	涿州、暴雨、救援、北京、物资、应急、救援队、影响、感谢、消防员、门头沟区、门头沟、河北、台风、房山区
3	同志通报追记	同志、救援、熊丽、北京市、群众、刘建民、牺牲、刘捷、房山、冯振、英勇、生前、北京、蓝天、王宏春
4	救援队伍与民警官兵赞扬	救援、致敬、冯振、群众、北京、英雄、涿州、被困、村民、副镇长、子弟兵、组织、消防员、一线、暴雨
5	技术运用助力灾害恢复	通信、北京、无人机、河北、物资、武警官兵、小时、门头沟、紧急、居民、洪水、连续、奋战、保障、道路
6	指挥中心救援指导	大街、消防、指挥、体育、冯振、救援、风雨、话务量、涉汛、中心、施工、平安、通行、班长、队员

3.1.3 提示与生成

根据提取的主题以及主题下的高频词项，挑选出每类信息下较具有代表性的微博文本作为人类生成的信息，以指导信息为例，研究所选择用作实验的人类生成的信息内容如下：

“7月29日起京津冀地区出现大范围强降雨天气。7月29日20时至8月1日17时，北京地区全市平均降雨量270.2毫米，城区平均降雨量242.7毫米，最大降雨出现在昌平王家园水库744.8毫米；最大雨强出现在丰台千灵山，31日10—11时降雨量为111.8毫米/时。预计8月1日夜间降雨趋势总体减弱，雨量分布不均，全市有小到中雨，东北部局地暴雨，伴有雷电活动。2日仍有雷阵雨天气，局地伴有大风、短时雨强较大。预计2日至5日白天最高气温32~33℃，湿度较大，体感闷热。#防范建议#1.降雨造成路面湿滑、能见度下降、部分低洼路段仍有积水，出行请注意安全。2.山区和浅山区山洪、滑坡、崩塌、泥石流等次生灾害风险极高，需重点关注山区易塌方路段的交通出行安全，不要前往山区、河道、地质灾害隐患区域。3.雷雨和大风时不要在高大建筑物、广告牌、临时搭建物或大树下方停留。4.雨后气温回升，湿度较大，闷热感增强，请注意做好防暑措施。”

接下来，根据信息所对应的主题以及相应的高频词项，通过 ChatGPT 生成由人工智能生成的信息。本文采用基于 GPT-3.5 版本的 ChatGPT 服务；相比于早期版本，该版本在训练技术和数据集上有显著的改进，能够以最小的特定于任务的微调来执行广泛的任務，仅依赖提示（Prompt）生成相关回应^[59]。在与 ChatGPT 对话时，根据 LDA 所提取的主题，输入第一次提示：“请以降水防护预警为主题，为我生成一条暴雨期间的信息”，第一次对话结果得到了一条初始的暴雨防范信息。与 ChatGPT 的对话中，提示词是非常重要的概念，用于引导或激发模型生成特定类型的语言成分^[60]，不同的提示会得到不同的输

出。Giray^[61]在研究中提到提示词的要素包括指令、提供背景知识的外部信息、输入数据以及输出格式。为了生成与所选人类生成内容差异最小的信息,根据具体的信息内容不断调整提示词,在接下来的对话中加入了由 LDA 提取结果得到的特定数据及词语,同时加入输出类型的要求。经过了三次的对话训练后,输入的提示为:“假设你是政务平台的应急信息发布者,请以降水防护预警为主题,为我生成一条暴雨期间的指导信息,可以参考的词语有:降雨、暴雨、天气、雨强、气象台、北京、地区、雷阵雨、高温、防范建议”。最终得到的由人工智能生成的信息内容如下:

“北京气象台发布降水防护预警,预计暴雨天气将持续影响该地区。预计降雨量在 80 毫米左右,降雨持续时间预计为 4 到 6 小时,雨强较大,可能伴有雷阵雨。同时,由于高温天气,暴雨期间还可能出现局部的地质灾害。

为应对暴雨期间的降水情况,气象台提出以下防范建议:市民应密切关注天气预报和预警信息,尽量减少外出活动,确保自身安全。在室内时,保持门窗关闭以避免雨水进入室内。如需外出,请选择安全通行路线,避免涉水区域。同时,出行时应注意交通安全,防止因道路湿滑导致的交通事故。若遇到强降雨和雷阵雨,请远离高海拔地区、河流及地势较低的地方,避免触电和溺水风险。此外,暴雨天气引发的高温 and 潮湿环境也会增加疾病传播的风险,市民应加强个人卫生防护,保持良好的室内通风,并注意及时食用烹饪熟透的食物以防食源性疾病。”

通过与 ChatGPT 的对话训练,在指导信息、调整信息、支持信息下都得到了人工智能生成的应急响应信息。需要说明的是,在灾害发生期间,被划分为调整信息的内容,大多数为实际发生的真实救援事例,不适合用人工智能生成的信息作为实验材料,最终研究选择以人类和人工智能生成的指导信息、支持信息两组信息作为实验材料。根据上文的解释,指导信息具有一定的指示性以及专业性,支持信息涉及情绪的安抚,这两类信息作为实验材料也检验了 ChatGPT 对于不同风格语言的生成能力。

3.2 测量工具

通过信息生成方式(人类 vs.人工智能)×信息类型(指导信息 vs.支持信息)的混合设计实验来研究公众对于不同信息的感知质量。为尽量避免其他无关因素的影响,本文将不同信息统一格式,利用 PS 技术将信息嵌入政务微博发布的内容界面,不告知被试者信息的生成方式,每一组实验材料除内容外其他部分完全相同。实验通过线上问卷的方式开展,被试者随机阅读其中一则信息后回答问题,问卷的发放与收集均通过线上进行。

结合本文研究的目标与特点,参考已有且经过验证的量表,研究设计了初步的问卷,问卷采用李克特 7 级量表,范围从“1=非常不同意”至“7=非常同意”。量表分为两部分,第一部分包括基本的人口统计学特征,第二部分测量被试者阅读信息后的感知质量。具体测量了来源可信度、内在信息质量、相关性、及时性以及感知信息价值五个方面。根据 Wang 和 Strong^[12]的描述,内在信息质量是一个多维的结构,包括准确性、客观性、时效性、相关性、可靠性等,表示数据本身的质量。在 Baabdullah^[13]提出的 AI 对话代理信息质量模型中将内在信息质量作为检验信息质量的一个重要维度。Zhou 等^[14]指出,衡量指导信息和调整信息的有效性,需要研究信息在多大程度上保护了利益相关者,且从受众者的角度来评估信息质量。他们通过控制信息呈现(文本 vs.信息图表)提出了生动性对于提升信息有效性的假设并且研究结果之一是发现更高的生动性会带来更强的感知信任。因此,本文最终确定从准确、客观、生动三个角度来设置题项测量内在信息质量。另外,结合应急响应信息的特点,单独测量被试者对于信息及时性、相关性的感知。感知价值是指用户对产品价值或服务的主观评价^[15],本文参考感知价值的概念对感知信息价值进行测量。具体的变量及题项见表 4。

表 4 变量设计

变量	题项	对应内容	参考文献
来源可信度	SC1	我认为该信息的来源具有一定的权威性	[16, 17]
	SC2	我认为提供该信息的来源值得信赖	
内在信息质量	IQ1	我认为该信息是准确的	[13, 18, 19]
	IQ2	我认为该信息内容观点客观公正	
	IQ3	我认为该信息所提供的内容是丰富的	
相关性	IR1	我认为该信息与当时情境下的灾害应对相关	[17, 20]
	IR2	我认为该信息对于灾害应对是合适的	
及时性	IT1	在当时的情境下，我认为该信息的提供是及时的	[20, 21]
	IT2	在当时的情境下，我认为该信息是实时的	
感知信息价值	PIV1	我认为该信息所提供的内容是有价值的	[13, 20, 22]
	PIV2	我认为该信息会让我更加了解暴雨灾害相关的情况	
	PIV3	我认为该信息值得我信任	
	PIV4	我认为该信息内容不极端化或情绪化	

3.3 样本选择与数据收集

通过 G*Power 进行事先 power 分析，基于 0.05 的显著性水平和 0.9 的 power 水平，期望效应量 f 为 0.25，得到的最小样本量为 171。在正式发放问卷前，先进行了小范围的问卷前测，共收集 51 份样本。正式发布后，经过筛选排除不符合样本要求等无效样本，最终得到有效样本 404 份，问卷的总体有效率为 83.1%。样本的基本统计信息见表 5。

表 5 样本基本统计信息

项目	类型	数量	占比
性别	男	231	57.2%
	女	173	42.8%
年龄	低于 18 岁	1	0.2%
	18~25 岁	72	17.8%
	26~30 岁	122	30.2%
	31~40 岁	158	39.1%
	41~50 岁	33	8.2%
	51~60 岁	17	4.2%
	超过 60 岁	1	0.2%
教育水平	初中及以下	2	0.5%
	高中/中专	13	3.2%
	大学专科	53	13.1%

续表

项目	类型	数量	占比
教育水平	大学本科	301	74.5%
	研究生及以上	35	8.7%
职业	学生	43	10.6%
	服务从业者	154	38.1%
	制造/建筑/电/采矿	188	46.5%
	农业/牧民/渔民	11	2.7%
	退休/待业	8	2.0%

4 结果分析

4.1 信度和效度检验

本文采用 SPSS 25.0 和 AMOS 28 对样本数据进行信效度分析。对 51 份前测问卷样本进行初步信效度检验，量表整体以及所有变量的克隆龙赫系数在 0.701~0.853，均大于 0.7，且 Bartlett 球形检验的伴随概率均小于 0.001，所有测度项的因子载荷均大于 0.5，说明问卷有较好的内部一致性和效度。正式发布后收集到的 404 份有效样本信效度检验见表 6。需要说明的是，参考相关研究，当 CR 高于 0.6 时，可以接受小于 0.5 但是大于 0.4 的 AVE 值，结构效度仍然足够^[62]，因此问卷的效度也可以被接受。另外，通过游程检验对样本数据进行随机性检验，被测变量显著性均大于 0.05，因此，样本数据满足随机性要求。

表 6 信效度结果

变量	题项	克隆巴赫系数	标准化的因子载荷	CR	AVE
来源可信度	SC1	0.747	0.803	0.7485	0.598
	SC2		0.743		
内在信息质量	IQ1	0.731	0.732	0.7393	0.486
	IQ2		0.707		
	IQ3		0.651		
相关性	IR1	0.714	0.738	0.7182	0.560
	IR2		0.759		
及时性	IT1	0.748	0.805	0.7545	0.606
	IT2		0.751		
感知信息价值	PIV1	0.746	0.656	0.7462	0.425
	PIV2		0.592		
	PIV3		0.717		
	PIV4		0.636		
整体		0.907			

4.2 操纵检验

在问卷前测时对信息类型进行了操纵检验，通过对两类信息及定义的随机展示检验信息类型这一变量的操纵是否有效。结果显示，在阅读指导信息的被试者（ $N=26$ ）中，对于该类型属于指导信息的认同度显著高于支持信息（ $M_{\text{指导}}=6.27$ ， $SD=0.778$ ， $M_{\text{支持}}=1.88$ ， $SD=0.833$ ， $t=19.465$ ， $p<0.001$ ）。在阅读支持信息的被试者（ $N=25$ ）中，对于该类型属于支持信息的认同度显著高于指导信息（ $M_{\text{指导}}=2.04$ ， $SD=0.871$ ， $M_{\text{支持}}=6.12$ ， $SD=0.927$ ， $t=16.208$ ， $p<0.001$ ）。因此，信息类型的操纵成功。

4.3 差异性分析

对数据进行正态性、方差齐性的检验后，通过单因素方差分析检验不同信息生成（人类 vs.人工智能）对来源可信度、内在信息质量、相关性、及时性以及感知信息价值进行差异性分析。方差结果显示，尽管人工智能生成的信息在各方面的感知得分均稍高于人类生成的信息(图 2)，但是这些差异均不显著，这表明公众在面对不同角色生成的应急信息时，没有比较明显的感知差别，这也印证了部分研究的结论，人类无法分辨人类与人工智能所生成的信息。

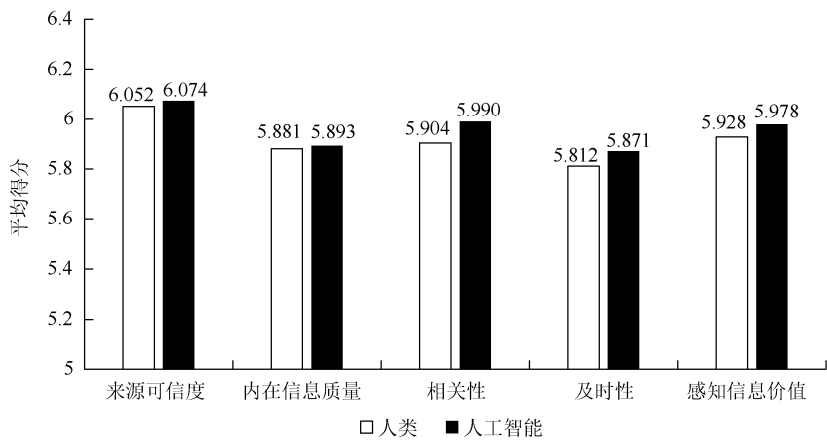


图 2 信息感知质量差异

接下来，针对信息生成方式（人类 vs.人工智能）和信息类型（指导信息 vs.支持信息）进行多因素方差分析，结果发现信息生成方式与信息类型之间不存在交互作用，信息类型对于测量变量影响显著（ $p=0.021$ ），且信息生成方式对所测量的变量影响仍不显著。具体地，如表 7 所示，研究发现指导信息在相关性、及时性、感知信息价值上得分均高于支持信息且结果显著，这一点印证了上文对于信息类型的不同解释。可以说，支持信息更加侧重情感层面的表达与鼓励，公众在阅读指导信息时可能感知到更直接的专业性与指导性。

表 7 信息类型的成对比较

变量	(I) 信息类型	(J) 信息类型	平均值差值 (I-J)	标准误差	显著性
来源可信度	指导信息	支持信息	0.131	0.074	0.077
	支持信息	指导信息	-0.131	0.074	0.077

续表

变量	(I) 信息类型	(J) 信息类型	平均值差值 (I-J)	标准误差	显著性
内在信息质量	指导信息	支持信息	0.134	0.072	0.064
	支持信息	指导信息	-0.134	0.072	0.064
相关性	指导信息	支持信息	0.191*	0.075	0.011
	支持信息	指导信息	-0.191*	0.075	0.011
及时性	指导信息	支持信息	0.183*	0.083	0.028
	支持信息	指导信息	-0.183*	0.083	0.028
感知信息价值	指导信息	支持信息	0.223*	0.064	0.001
	支持信息	指导信息	-0.223*	0.064	0.001

*的显著性水平为 0.05，**的显著性水平为 0.01

基于此，考虑到不同类型信息所传达内容以及情感的不同，为避免其中的影响，研究进一步单独对指导信息和支持信息进行单因素方差分析，结果仍然均不显著（表 8），再一次说明公众对于人类和人工智能生成的信息没有较为优异的区分能力。

表 8 信息分类的单因素方差分析

变量	生成方式	指导信息 (N=202)				支持信息 (N=202)			
		M	SD	F	p	M	SD	F	p
来源可信度	人类 (N=101)	6.134	0.651	0.011	0.918	5.970	0.806	0.236	0.628
	人工智能 (N=101)	6.124	0.716			6.025	0.789		
内在信息质量	人类 (N=101)	5.980	0.671	0.320	0.572	5.782	0.755	0.482	0.489
	人工智能 (N=101)	5.927	0.654			5.858	0.799		
相关性	人类 (N=101)	6.045	0.693	0.003	0.959	5.762	0.777	2.459	0.118
	人工智能 (N=101)	6.040	0.688			5.941	0.838		
及时性	人类 (N=101)	5.965	0.722	0.389	0.534	5.658	0.930	1.971	0.162
	人工智能 (N=101)	5.901	0.745			5.842	0.924		
感知信息价值	人类 (N=101)	6.102	0.640	0.750	0.388	5.755	0.700	3.358	0.068
	人工智能 (N=101)	6.027	0.578			5.928	0.643		

为了探讨影响人们感知信息质量的因素,基于人口统计学的特征,单独对指导信息和支持信息的生成方式和性别、年龄以及教育水平进行差异性分析。首先,所有情况下,信息的两种生成方式对于感知信息质量的不同维度影响不显著,关于这一点,可能与人们目前缺乏对于网络信息的系统式思考有关^[63]。性别的显著性差异体现在支持信息对于内在信息质量的感知中,结果显示女性的平均得分要高于男性($M_{女}-M_{男}=0.232, p=0.037$),该结果说明女性在支持信息中感受到了更高的信息质量,这一点可能与男女之间的感性与理性的思维差异相联系,而在指导信息中,性别没有显著性差异。

在年龄和教育水平方面,显著性差异出现在支持信息中,研究结果表现出一个有趣的现象,在信息来源可信度的感知上,高年龄群体(60岁以上)的得分显著低于青年(26~30岁)和中年(31~40岁)这两个群体,这说明老年群体对于表达情感支持的信息的接收表现出更为谨慎的态度。

值得一提的是,本文在指导信息中发现了教育水平的显著效应。具体而言,指导信息下的显著性差异发生在中等学历之间。在指导信息及时性的感知上,本科和研究生及以上学历存在显著性差异,这说明对于高学历群体来说,他们对于信息时效的要求会更高。在内在信息质量和相关性上,显著性差异出现在专科组和本科组之间,但并不是专科组的平均得分高于本科组,而是专科组的平均得分略低于本科组,这说明对于信息质量的感知也存在着例外,并不是都随着学历的升高而更加严苛(见附录)。

5 结论与展望

5.1 研究结论

本文设计了一个基于大模型和社交媒体数据的应急信息生成框架,基于情境危机沟通理论研究三类应急信息,应用LDA模型提取主题和关键词项优化大模型的提示词,从而生成应急响应信息,探讨了人工智能生成应急信息的感知质量以及相较于人类生成是否存在差异,进一步研究这种差异在不同人群中的具体表现。主要结论可以总结为两方面:

首先,针对研究问题①,本文研究发现,使用ChatGPT基于微博数据生成的应急信息在公众中的感知质量与人类生成的应急信息并无显著区别。这一发现与其他领域^[8, 9]关于人工智能生成信息的研究结论基本一致,显示出人工智能在生成高质量应急信息方面的巨大潜力。并且在感知内容质量的主要维度上,人工智能生成的应急信息都有较稳定的得分表现,表明生成式人工智能技术在应急信息生成方面的应用可行性。另外,与一项采用大模型生成公共卫生信息的研究^[4]相比,本文将情境危机沟通理论、主题提取模型等理论和技术融合到大模型生成应急信息过程中,这有助于提升信息生成质量。

进一步地,针对研究问题②,本文研究揭示了针对不同危机应对策略,不同信息主体的功能特征,以及基于人口统计学特征的性别、年龄和教育水平对应急响应感知信息质量的影响存在显著差异。具体而言:性别对支持信息的感知内在信息质量存在显著性差别,表明不同性别在处理信息时的偏好和敏感度可能不同;年龄对信息来源可信度有着显著影响,结果显示年老者对于某些信息持有更加谨慎的态度,这可能与其对信息真实性的更高要求有关;教育水平对指导信息的感知信息价值、相关性以及及时性均有着显著的感知差别,说明教育背景对信息接收和解读的能力有重要影响。

这些结论不仅验证了生成式人工智能在应急信息生成中的有效性,还揭示了不同群体在信息质量感知上的差异性,为未来的危机信息管理策略提供了重要参考依据。

5.2 研究启示

本文的理论贡献主要有三方面。首先,结合情境危机沟通理论,设计了针对不同类型应急信息的大模型提示词优化方法,为应急信息生成中大模型的应用范式提供了有益参考。无论是危机初期的紧急指导,还是后期的情绪安抚和信任重建,基于情境危机沟通理论的分类优化方法能够快速调整和优化生成内容,更好地满足公众的动态需求。其次,从信息质量的角度探讨了生成式人工智能技术在应急信息生成中的应用可行性,发现了人口统计学特征对人工智能生成应急信息感知质量的影响作用,为优化应急信息生成策略提供理论依据。与诗歌创作^[8]、产品文案撰写^[9]等应用场景相比,本文补充了应急场景下用户信息使用行为的相关空白;进一步,人口统计学特征影响人工智能生成信息感知质量这一发现不仅适用于应急信息管理等领域,也对产品营销等领域有一定启示。最后,将情境危机沟通等理论与主题提取、生成式人工智能等技术相结合,扩大了相关理论在新一代人工智能技术背景下的适用范围,拓展了对危机信息管理、应急信息传播等领域下人智协作范式的认识。国内外学者呼吁重视 AIGC、大模型等新技术应用过程中的人智协作研究^[37, 45],在本文采用的应急信息生成框架中,没有单纯地依赖大模型技术,而是融合了经典理论对不同类型的应急信息分开考虑,并且研究人员和大模型进行了多轮对话,从而达到较好的生成效果,这在一定程度上体现了人智协作的作用。

本文对于将大模型等生成式人工智能技术应用到危机信息管理的实践活动具有一定的参考价值。首先,在进行生成式人工智能普及的过程中,相关人员可以针对重点人群进行有针对性的科普宣传。其次,政务社交媒体运营管理人员、危机信息管理的参与者等角色在应用生成式人工智能技术的过程中,应从突发事件的类型等角度出发分析受影响人群的构成情况,充分考虑人口统计学特征对应急信息感知质量评估等方面的影响,设计有针对性的危机信息管理策略,提升社交媒体应急信息对公众进行建议与安抚的实际效果。再次,本文所设计的应急信息生成框架能够为基于大模型的应急信息生成与管理系统提供系统设计方面的参考。最后,本文研究结合社交媒体数据并考虑应急信息的类型划分;对此,建议相关部门利用社交媒体历史数据建立一个高质量的应急信息知识库。通过使用该知识库,相关人员能够参考过往事件,按照突发事件的不同类型,根据所需要发布应急信息的类型,检索相关的关键词项,对大模型的提示词进行优化,从而提高大模型生成应急信息的效率。

5.3 研究局限和未来研究

本文也存在一定的局限。人类没有较好地区分人类和人工智能生成的信息,但这并不能说明人工智能生成技术已经成熟,在不同领域关于人工智能生成的伦理问题也一直在被讨论^[64],生成式人工智能技术在危机信息管理中的应用仍存在一定的挑战,如大模型的“幻觉”问题^[45]。本次研究所生成的信息避开了较多客观事例以及数据内容,避免了 ChatGPT 编造故事和数据的可能,用作实验的文本材料侧重于传达防范建议和情感鼓励,在实际应用过程中,实时更新的数据的准确性以及内容的真实性仍是不容忽视的重要问题。此次研究的目的在于初步探索人类在突发危机时对于人工智能生成信息的感知,并挖掘人口统计学特征对于信息感知差异的影响,为政府、媒体机构等应急信息传播主体提供相关建议。未来希望有机会能够进一步将实验放在更加真实的环境中或收集实证数据,从而更加准确地探究人类对人工智能生成应急信息的感知质量。

参考文献

- [1] 周利敏, 钟娇文. 应急管理中社交媒体的嵌入: 理论构建与实践创新[J]. 中国行政管理, 2022, 38 (1): 121-127.
- [2] 孙宗锋, 郑跃平. 我国城市政务微博发展及影响因素探究: 基于 228 个城市的“大数据+小数据”分析(2011—2017)[J]. 公共管理学报, 2021 (1): 77-89, 171.
- [3] 田沛霖, 赵蓉英. 国际应急信息管理研究进展与趋势[J]. 图书馆学研究, 2022, (10): 15-23.
- [4] Karinshak E, Liu S X, Park J S, et al. Working with AI to persuade: examining a large language model's ability to generate pro-vaccination messages[J]. Proceedings of the ACM on Human-Computer Interaction, 2023, 7 (CSCW1): 1-29.
- [5] Bommasani R, Hudson D A, Adeli E, et al. On the opportunities and risks of foundation models[J]. arxiv preprint arxiv: 2108.07258, 2021.
- [6] Chandra A, Chakraborty A. Exploring the role of large language models in radiation emergency response[J]. Journal of Radiological Protection, 2024, 44 (1): 011510.
- [7] Pereira J, Fidalgo R, Lotufo R, et al. Crisis event social media summarization with GPT-3 and neural reranking[C]//Proceedings of the 20th International ISCRAM Conference. Omaha: University of Nebraska at Omaha, 2023: 371-384.
- [8] Köbis N, Mossink L D. Artificial intelligence versus Maya Angelou: experimental evidence that people cannot differentiate AI-generated from human-written poetry[J]. Computers in Human Behavior, 2021, 114: 106553.
- [9] Zhang Y H, Gosline R. Human favoritism, not AI aversion: people's perceptions (and bias) toward generative AI, human experts, and human-GAI collaboration in persuasive content generation[J]. Judgment and Decision Making, 2023, 18: e41.
- [10] 王羽西, 于兴尚, 徐中阳. 基于公众感知的政府应急网站信息服务质量可拓性评价研究[J]. 情报科学, 2024, 42 (9): 155-164, 177.
- [11] 王敏, 陈娟. 应急沟通共同体: 面向复杂风险治理的新议题[J]. 中国行政管理, 2024, 40 (12): 141-152.
- [12] Wang R Y, Strong D M. Beyond accuracy: what data quality means to data consumers[J]. Journal of Management Information Systems, 1996, 12 (4): 5-33.
- [13] Baabdullah A M. Generative conversational AI agent for managerial practices: the role of IQ dimensions, novelty seeking and ethical concerns[J]. Technological Forecasting and Social Change, 2024, 198: 122951.
- [14] Zhou Z Y, Zhang X Y, Ki E J. Enhancing the efficacy of instructing and adjusting information: an exploration of interactivity and vividness[J]. Journal of Contingencies and Crisis Management, 2023, 31 (4): 809-825.
- [15] 周涛, 李松洮, 邓胜利. 用户信息搜寻转移意向研究: 从搜索引擎到生成式 AI[J]. 图书情报工作, 2024, 68 (3): 49-58.
- [16] 岑咏华, 王晓书, 万青, 等. 个体信息认知处理与态度形成机制的实证研究[J]. 管理学报, 2016, 13 (6): 880-888.
- [17] Bhattacharjee A, Sanford C. Influence processes for information technology acceptance: an elaboration likelihood model[J]. MIS Quarterly, 2006, 30 (4): 805-825.
- [18] Jiang Z H, Benbasat I. The effects of interactivity and vividness of functional control in changing web consumers' attitudes[C]//ICIS 2003 Proceedings. New York: AIS, 2003: 93.
- [19] Kelley C A, Gaidis W C, Reingen P H. The use of vivid stimuli to enhance comprehension of the content of product warning messages[J]. Journal of Consumer Affairs, 1989, 23 (2): 243-266.
- [20] Cheung C M K, Lee M K O, Rabjohn N. The impact of electronic word-of-mouth: the adoption of online opinions in online customer communities[J]. Internet Research, 2008, 18 (3): 229-247.
- [21] 杨雪梅. 科研社交网络环境下信息分享行为影响因素研究[D]. 武汉: 武汉大学, 2019.
- [22] Sussman S W, Siegal W S. Informational influence in organizations: an integrated approach to knowledge adoption[J]. Information Systems Research, 2003, 14 (1): 47-65.
- [23] Moravec V, Hynek N, Skare M, et al. Human or machine? The perception of artificial intelligence in journalism, its

- socio-economic conditions, and technological developments toward the digital future[J]. *Technological Forecasting and Social Change*, 2024, 200: 123162.
- [24] Xu J P, Wang Z Q, Shen F, et al. Natural disasters and social conflict: a systematic literature review[J]. *International Journal of Disaster Risk Reduction*, 2016, 17: 38-48.
- [25] 丁学君, 樊荣, 苗蕊, 等. 社会化媒体中突发公共卫生事件舆情传播规律研究[J]. *信息系统学报*, 2018, (2): 1-14.
- [26] 章钢, 谢阳群. 危机信息管理研究综述[J]. *情报杂志*, 2006, 25 (8): 22-25.
- [27] Wang Y D, Wang T, Ye X Y, et al. Using social media for emergency response and urban sustainability: a case study of the 2012 Beijing rainstorm[J]. *Sustainability*, 2015, 8 (1): 25.
- [28] Li L F, Zhang Q P, Tian J, et al. Characterizing information propagation patterns in emergencies: a case study with Yiliang Earthquake[J]. *International Journal of Information Management*, 2018, 38 (1): 34-41.
- [29] Ogie R I, James S, Moore A, et al. Social media use in disaster recovery: a systematic literature review[J]. *International Journal of Disaster Risk Reduction*, 2022, 70: 102783.
- [30] 张渝, 黄慧敏. 政府在线回应质量感知对公众政治信任与持续电子参与意愿的影响: 混合研究视角[J]. *国际新闻界*, 2024, 46 (7): 6-27.
- [31] 毛太田, 汤滢, 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.
- [32] 李白杨, 白云, 詹希旎, 等. 人工智能生成内容 (AIGC) 的技术特征与形态演进[J]. *图书情报知识*, 2023, 40 (1): 66-74.
- [33] Otal H T, Abdullah Canbaz M. AI-powered crisis response: streamlining emergency management with LLMs [C]//2024 IEEE World Forum on Public Safety Technology (WFPST). Herndon: IEEE, 2024: 104-107.
- [34] 张夏恒, 马妍. AIGC 在应急情报服务中的应用研究[J]. *图书馆工作与研究*, 2024, (11): 60-67.
- [35] Colverd G, Darm P, Silverberg L, et al. FloodBrain: flood disaster reporting by web-based retrieval augmented generation with an LLM[J]. *arXiv preprint arXiv: 2311.02597*, 2023.
- [36] Otal H T, Stern E, Canbaz M A. LLM-assisted crisis management: building advanced LLM platforms for effective emergency response and public collaboration [C]//2024 IEEE Conference on Artificial Intelligence (CAI). Singapore: IEEE, 2024: 851-859.
- [37] 沙勇忠, 何路, 李潇敏. 公共危机信息管理研究进展与趋势[J]. *情报学进展*, 2024, 15 (00): 84-130.
- [38] Maltz E. Is all communication created equal? : an investigation into the effects of communication mode on perceived information quality[J]. *Journal of Product Innovation Management*, 2000, 17 (2): 110-127.
- [39] van Dalen A. The algorithms behind the headlines: how machine-written news redefines the core skills of human journalists[J]. *Journalism Practice*, 2012, 6 (5/6): 648-658.
- [40] 王平, 程齐凯. 网络信息可信度评估的研究进展及述评[J]. *信息资源管理学报*, 2013, 3 (1): 46-52.
- [41] 苏君华, 杜念. 国外应急信息质量评价研究进展与趋势思考[J]. *图书情报知识*, 2024, 41 (1): 143-154.
- [42] Seppänen H, Virrantaus K. Shared situational awareness and information quality in disaster management[J]. *Safety Science*, 2015, 77: 112-122.
- [43] Zhou J W, Zhang Y X, Luo Q N, et al. Synthetic lies: understanding AI-generated misinformation and evaluating algorithmic and human solutions[C]//Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems. New York: ACM, 2023: 1-20.
- [44] Fügener A, Grahl J, Gupta A, et al. Will humans-in-the-loop become borgs? Merits and pitfalls of working with AI[J]. *MIS Quarterly*, 2021, 45 (3): 1527-1556.
- [45] Dwivedi Y K, Kshetri N, Hughes L, et al. Opinion Paper: “So what if ChatGPT wrote it?” Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy[J]. *International Journal of Information Management*, 2023, 71: 102642.
- [46] Coombs W T. Choosing the right words: the development of guidelines for the selection of the “appropriate” crisis-response strategies[J]. *Management Communication Quarterly*, 1995, 8 (4): 447-476.

- [47] Liu W L, Lai C H, Xu W A. Tweeting about emergency: a semantic network analysis of government organizations' social media messaging during Hurricane Harvey[J]. *Public Relations Review*, 2018, 44 (5): 807-819.
- [48] Coombs W T. Protecting organization reputations during a crisis: the development and application of situational crisis communication theory[J]. *Corporate Reputation Review*, 2007, 10 (3): 163-176.
- [49] Li Y R, Chandra Y, Fan Y Y. Unpacking government social media messaging strategies during the COVID-19 pandemic in China[J]. *Policy & Internet*, 2022, 14 (3): 651-672.
- [50] 胡峰. 活动理论视阈下公共卫生应急信息公开的运行结构透视: 要素解析与系统刻画[J]. *情报资料工作*, 2024, 45 (6): 40-49.
- [51] Kim S, Liu B F. Are all crises opportunities? A comparison of how corporate and government organizations responded to the 2009 flu pandemic[J]. *Journal of Public Relations Research*, 2012, 24 (1): 69-85.
- [52] Page T G. The reputational benefits of instructing information: the first test of the revised model of reputation repair[J]. *Public Relations Review*, 2022, 48 (5): 102256.
- [53] Hughes A L, St Denis L A A, Palen L, et al. Online public communications by police & fire services during the 2012 Hurricane Sandy[C]//*Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*. New York: ACM, 2014: 1505-1514.
- [54] Zheng L F, Chen L, Long F J, et al. Reducing social media attention inequality in disasters: the role of official media during rainstorm disasters in China[J]. *International Journal of Disaster Risk Science*, 2024, 15 (3): 388-403.
- [55] Page T G. Measuring success: explications and measurement scales of instructing information and adjusting information[J]. *Public Relations Review*, 2020, 46 (4): 101952.
- [56] 陈晓美, 高铨, 关心惠. 网络舆情观点提取的 LDA 主题模型方法[J]. *图书情报工作*, 2015, 59 (21): 21-26.
- [57] 赵天资, 段亮, 岳昆, 等. 基于 Biterm 主题模型的新闻线索生成方法[J]. *数据分析与知识发现*, 2021, 5 (2): 1-13.
- [58] 池毛毛, 潘美钰, 王伟军. 共享住宿与酒店用户评论文本的跨平台比较研究: 基于 LDA 的主题社会网络和情感分析[J]. *图书情报工作*, 2021, 65 (2): 107-116.
- [59] Talebi S, Tong E, Mofrad M R K. Beyond the hype: assessing the performance, trustworthiness, and clinical suitability of GPT3.5[J]. *arxiv preprint arxiv: 2306.15887*, 2023.
- [60] 陈秋心, 邱泽奇. “人机互生”时代可供性理论的契机与危机: 基于“提示词”现象的考察[J]. *苏州大学学报(哲学社会科学版)*, 2023, 44 (5): 172-182.
- [61] Giray L. Prompt engineering with ChatGPT: a guide for academic writers[J]. *Annals of Biomedical Engineering*, 2023, 51 (12): 2629-2633.
- [62] Fan L, Wang Y W, Mou J. Enjoy to read and enjoy to shop: an investigation on the impact of product information presentation on purchase intention in digital content marketing[J]. *Journal of Retailing and Consumer Services*, 2024, 76: 103594.
- [63] Jakesch M, Hancock J T, Naaman M. Human heuristics for AI-generated language are flawed[J]. *Proceedings of the National Academy of Sciences*, 2023, 120 (11): e2208839120.
- [64] Guo D H, Chen H X, Wu R L, et al. AIGC challenges and opportunities related to public safety: a case study of ChatGPT[J]. *Journal of Safety Science and Resilience*, 2023, 4 (4): 329-339.

附录 教育水平的多重比较

变量	(I) 教育水平	(J) 教育水平	平均值差值 (I-J)	标准误差	显著性
内在信息质量	初中及以下	高中/中专	0.194	0.503	0.699
		大学专科	0.299	0.481	0.536
		大学本科	-0.028	0.469	0.952
		研究生及以上	0.146	0.494	0.768
	高中/中专	初中及以下	-0.194	0.503	0.699
		大学专科	0.104	0.226	0.645
		大学本科	-0.222	0.198	0.263
		研究生及以上	-0.049	0.251	0.847
	大学专科	初中及以下	-0.299	0.481	0.536
		高中/中专	-0.104	0.226	0.645
		大学本科	-0.327 [*]	0.134	0.016
		研究生及以上	-0.153	0.205	0.457
	大学本科	初中及以下	0.028	0.469	0.952
		高中/中专	0.222	0.198	0.263
		大学专科	0.327 [*]	0.134	0.016
		研究生及以上	0.174	0.174	0.318
	研究生及以上	初中及以下	-0.146	0.494	0.768
		高中/中专	0.049	0.251	0.847
		大学专科	0.153	0.205	0.457
		大学本科	-0.174	0.174	0.318
相关性	初中及以下	高中/中专	0.958	0.519	0.066
		大学专科	0.957	0.497	0.056
		大学本科	0.624	0.484	0.199
		研究生及以上	0.906	0.510	0.077
	高中/中专	初中及以下	-0.958	0.519	0.066
		大学专科	-0.001	0.233	0.995
		大学本科	-0.334	0.204	0.103
		研究生及以上	-0.052	0.259	0.841
	大学专科	初中及以下	-0.957	0.497	0.056
		高中/中专	0.001	0.233	0.995

续表

变量	(I) 教育水平	(J) 教育水平	平均值差值 (I-J)	标准误差	显著性
相关性	大学专科	大学本科	-0.333 [*]	0.138	0.017
		研究生及以上	-0.051	0.212	0.811
	大学本科	初中及以下	-0.624	0.484	0.199
		高中/中专	0.334	0.204	0.103
		大学专科	0.333 [*]	0.138	0.017
		研究生及以上	0.282	0.179	0.117
相关性	研究生及以上	初中及以下	-0.906	0.510	0.077
		高中/中专	0.052	0.259	0.841
		大学专科	0.051	0.212	0.811
		大学本科	-0.282	0.179	0.117
及时性	初中及以下	高中/中专	-0.250	0.558	0.655
		大学专科	-0.147	0.534	0.784
		大学本科	-0.236	0.520	0.650
		研究生及以上	0.250	0.548	0.649
	高中/中专	初中及以下	0.250	0.558	0.655
		大学专科	0.103	0.251	0.680
		大学本科	0.014	0.219	0.949
		研究生及以上	0.500	0.279	0.075
	大学专科	初中及以下	0.147	0.534	0.784
		高中/中专	-0.103	0.251	0.680
		大学本科	-0.090	0.149	0.548
		研究生及以上	0.397	0.227	0.083
	大学本科	初中及以下	0.236	0.520	0.650
		高中/中专	-0.014	0.219	0.949
		大学专科	0.090	0.149	0.548
		研究生及以上	0.486 [*]	0.193	0.012
	研究生及以上	初中及以下	-0.250	0.548	0.649
		高中/中专	-0.500	0.279	0.075
		大学专科	-0.397	0.227	0.083
		大学本科	-0.486 [*]	0.193	0.012

*显著性水平为 0.05

Research on Perceived Quality of Emergency Information Generated by Artificial Intelligence based on Government Weibo Data

LI Yixiao¹, ZHANG Lu¹, ZHOU Shasha², YANG Shuiqing²

(1. School of Information Technology and Artificial Intelligence, Zhejiang University of Finance & Economics, Hangzhou 310018, China;

2. School of Management, Zhejiang University of Finance & Economics, Hangzhou 310018, China)

Abstract The immediate production and dissemination of emergency information is very important for emergency response. In the context of its widespread application, generative artificial intelligence technologies and tools can improve the efficiency of producing emergency information. This paper crawls government Weibo data released during a rainstorm disaster and classifies the texts into three categories according to the situational crisis communication theory. The LDA model is adopted to extract text topics and core terms, which are further used to optimize prompts for large language models to generate targeted emergency response information. This paper further explores public perceptual quality differences between human-written and AI-generated emergency information, as well as the influences of demographic characteristics on such discrepancies. The results show that there is no significant difference in the perceived quality of information generated by humans and artificial intelligence, but factors such as education and age significantly affect the perceived quality of different types of information. This paper discusses the feasibility of applying artificial intelligence generation technology and tools to the process of emergency information generation and provides relevant suggestions for the government, media organizations, and other emergency information dissemination bodies.

Key words AIGC, Emergency information, LDA, Perceived quality, ChatGPT

作者简介

李一啸 (1982—), 男, 浙江财经大学信息技术与人工智能学院教授、硕士生导师, 研究方向包括危机信息管理、网络信息管理、复杂系统等。E-mail: yxli@zufe.edu.cn。

张璐 (2000—), 女, 浙江财经大学信息技术与人工智能学院 2023 级硕士研究生, 研究方向为危机信息管理。E-mail: luz_1110@163.com。

周莎莎 (1987—), 女, 浙江财经大学管理学院副教授、硕士生导师, 研究方向包括在线消费者行为等。E-mail: zss_1224@163.com。

杨水清 (1977—), 男, 浙江财经大学管理学院教授、博士生导师, 研究方向包括人工智能与管理、数据要素与企业绩效等。E-mail: yangshuiqing@zufe.edu.cn。