

生成式智能搜索结果可信度研究

吴丹 孙国烨

摘要 以 ChatGPT 为代表的生成式人工智能取得突破性进展,使生成式智能搜索在运行机制和结果呈现形式上与传统搜索引擎的差异更加显著,也使得生成式智能搜索这一新型搜索工具的结果可信度面临新的挑战。本研究以 SOR 模型为基础,基于可信度研究的成果和想象可供性理论,从用户信息获取的视角出发,运用卡片分类法和叙事研究法,通过生成式智能搜索典型代表 New Bing 开展实验,旨在厘清生成式智能搜索结果可信度从刺激、机体到反应的个体心理链路,探究其形成机理与影响因素。以生成式智能搜索结果界面的信息线索为外部刺激,以用户的内在感知和信任意愿为机体,将导致的用户情绪及行为作为反应,最终形成了生成式智能搜索可信度的 SOR 模型。对生成式智能搜索可信度的形成机理进行分析后认为,系统的研发者可以面向用户,根据使用体验催生的搜索需求,从内容整合性、问答交互性、工具拟人性三个维度出发,优化生成式智能搜索系统。图 4。表 2。参考文献 54。

关键词 生成式人工智能 信息检索 可信度 Microsoft Bing

分类号 G254.9

The Credibility of the Results of Generative Intelligent Search

WU Dan, SUN Guoye

ABSTRACT

Along with the breakthrough of generative AI represented by ChatGPT, generative intelligent search is more significantly different from traditional search engines in operating mechanism and result presentation form, which makes the credibility of the results of this new search tool face new challenges. Since users' trust and willingness to accept generative intelligent search is particularly important, based on SOR model, combined with the results of credibility research and imagined affordances theory, this study focuses on the users, and makes up for the research gap on the credibility of generative intelligent search results from the perspective of individual information acquisition.

This study adopts card sorting and narrative research methods, and conducts experiments through New Bing, a typical representative of generative intelligent search, aiming to clarify the individual psychological link of the credibility of generative intelligent search results from stimulus, organism to response, and explore its formation mechanism and influencing factors.

First of all, this study summarizes three characteristic dimensions of generative intelligent search, which are content integration, question-and-answer interactivity, and tool imitation humanity.

Secondly, this study takes the information clues in the interface of generative intelligent search results as external stimuli, takes the users' internal perception and trust willingness in three characteristic dimensions

通信作者:吴丹,Email:woodan@whu.edu.cn,ORCID:0000-0002-2611-7317 (Correspondence should be addressed to WU Dan,Email:woodan@whu.edu.cn,ORCID:0000-0002-2611-7317)

as the organism, and takes the users' emotions and behaviors caused by credibility changes as responses, to form the SOR analytical framework of generative intelligent search credibility.

Subsequently, this study divided the elements of the search result interface of New Bing into seven types of information cue stimulus labels. Participants were asked to classify the cards corresponding to the clues into different trust perceptions, and answered which type of characteristic perception was stimulated by these clues, as well as whether they trusted New Bing or traditional search engines more in three characteristic perception dimensions. Participants were further asked to rank the degree to which three characteristic perceptions affected the credibility of search results, and to indicate how they would react when the credibility of a characteristic perception dimension was reduced.

Through research data analysis, this study identifies the information clues that have an impact on the credibility of search results, and which information cues stimulate each of the three characteristic perceptions. In addition, this study extracts the perception attributes corresponding to the three characteristic perception, and the user response caused by the reduction of credibility under each characteristic perception dimension. Finally, this study constructs a complete SOR model of generative intelligent search credibility. After a detailed analysis of the formation mechanism of generative intelligent search credibility, this study provides specific suggestions for system developers to optimize generative intelligent search system from three characteristic perception dimensions. 4 figs. 2 tabs. 54 refs.

KEY WORDS

Generative AI. Information retrieval. Credibility. Microsoft Bing.

0 引言

2022年底, OpenAI 公司研发了一款基于 GPT-3.5 大语言模型的聊天机器人 ChatGPT, 引发了关于生成式人工智能的全球热议。2023 年 2 月, Microsoft 发布的新版 Bing (以下称 New Bing) 整合了 GPT 背后的人工智能技术, 用户可以选择通过搜索框与 New Bing 进行对话, 以获取问题的针对性回答, 这一新兴技术与传统搜索引擎的结合更是让关于生成式人工智能的讨论达到高潮。诸如 OpenAI 的 ChatGPT、百度的文心一言、谷歌的 Bard 等聊天机器人目前仍停留在对话交流层面, 而 New Bing 则将相关技术集成到传统搜索引擎中, 标注回答的信息参考来源, 支持与网页搜索结果相结合进行答案溯源, 成为当前生成式智能搜索的典型代表。

基于大语言模型的生成式智能搜索在人工智能和信息检索领域都是具有变革性意义的工

具, 但其存在一个关键且无法回避的问题——可信度。可信度是用户检索信息时最为关注的维度之一, 生成式智能搜索的可信度问题也引发了各界的深切担忧。Nature 于 New Bing 问世的同月发布评论性文章, 认为人工智能聊天机器人被引入搜索引擎, 带来了搜索的新状态, 这将改变用户与搜索工具的关系, 这种形式的人机交互可能存在风险, 主要体现在结果的可信度存疑。受制于训练语料、成本等因素, 生成式智能搜索常常由于不能提取相关信息, 或无法区分可信和不可信的来源, 而为它们不知道答案的问题编造虚构的答案, 或自信地回答错误的信息, 这种现象被人工智能领域的学者称为“幻觉”^[1]。在仅使用生成式智能搜索的情况下, 用户难以筛选和验证结果, 它看似有理有据的推理逻辑, 能够将虚假和有误导性的内容捏造成令人信服的回答, 利用光环效应使用户受骗。

生成式智能搜索结果可信度过低会带来一

系列负面影响。对于个体用户而言,用户在发现算法的瑕疵后往往会拒绝使用它们^[2],进而对基于聊天的搜索类工具失去信心。对于社会而言,一方面,如果生成式智能搜索返回错误或具有偏见的信息,可能会导致误导性信息的进一步传播,但人工智能并不会对此负责任,反而总是以服从命令为由来开脱;另一方面,社会对工具的整体信任度降低,会导致新兴技术的发展受到阻碍。

在生成式人工智能的影响力逐步扩大、应用领域更加广泛的背景下,全球范围内对可信人工智能的呼吁也更为迫切。在法规准则层面,对生成式人工智能可信度的关注度极高。早在2019年4月,欧盟委员会就发布了《可信人工智能的伦理准则》,提出实现可信人工智能应遵循的核心原则和价值观。自生成式人工智能获得高度关注后,更多相关文件陆续出台,以推动发展可信人工智能。2023年6月14日,欧盟通过了全球第一部综合性人工智能法规——《人工智能法案》,旨在确保人工智能的透明度、责任和可控性。2023年7月13日,我国网信办等七部门联合公布《生成式人工智能服务管理暂行办法》,针对传播虚假信息等问题,统筹管理生成式人工智能发展和应用。在学术界,已有学者开展生成式人工智能内容评测^[3]、挑战应对^[4]、治理路径^[5]、监管策略^[6]等方面的研究,提出人工智能赋能的替代信息搜索理论框架^[7]。但现有研究均着眼于 ChatGPT,而未聚焦用户,缺少从个体信息获取的视角出发,对侧重于集成检索的生成式智能搜索可信度的系统研究。

用户与搜索工具交互的结果页面是信息检索研究领域的重要分支,也影响着用户的信息行为^[8]。生成式智能搜索在隐性的运行机制和显性的搜索结果界面呈现上显示出与传统搜索引擎的显著差异。生成式智能搜索支持以结构性文本的形式生成答案集合,呈现出语义要素之间的价值关联;可以与用户进行多轮对话,根据检索语境和意图,提供更贴合对话上下文和

用户需求的答案;倾向于在对话语句中加入拟人化和情感化的表达。

那么,这些运行机制和结果界面的显著差异会对生成式智能搜索结果的可信度产生何种影响?用户在使用生成式智能搜索时,通常认为哪些具体的结果界面可见元素是与可信度相关的?这些元素会使用户感知到信任方面的哪些具体属性,进而影响到他们对搜索结果的信任程度?信任程度的变化又会进一步导致用户的使用行为发生怎样的改变?为了解决以上问题,本研究从生成式智能搜索结果界面的信息线索出发,聚焦其对生成式智能搜索特征感知的影响,探讨生成式智能搜索结果可信度的形成机理,以及可信度变化导致的行为反应,为进一步优化生成式智能搜索提供参考。

1 文献回顾与理论基础

1.1 可信度研究概述

关于可信度的概念学界尚未形成统一的定义,但关于其内涵存在着共识,即可信度是从用户角度出发所感知到的准确性。可信度一词通常与信息质量挂钩,但可信度并不等同于信息质量本身,而是由用户多维度的主观感知所决定^[9]。总的来说,可信度是一个复合概念,它最为关键的内涵是值得被用户信赖。在图书情报领域,可信度主要应用于信息检索研究中,它能够影响用户对信息的接收,因而被作为信息检索效果的判断依据^[10]。

可信度的形成机理是信息检索领域的重要研究内容,学界一直以来对其予以高度关注,并形成了一系列经典理论。Pirolli 和 Card 提出的信息觅食理论将用户获取信息的行为比作动物觅食行为,认为用户根据信息线索来搜寻信息,与可信度相关的线索会影响用户的信息获取^[11]。Fogg 提出的显著性—阐释理论认为,用户首先在可信度评估过程中注意到显著性线索,随后基于对线索的认知和理解进行判断^[12]。Flanagin 和 Metzger 提出的双处理模型认为,用

户评估信息可信度的路径有两种,一种是启发式处理,通常根据少量的信息线索形成判断,另一种是系统式处理,通过对所有信息线索的综合考量,得出最终的判断^[13]。Sundar 提出的 MAIN 模型认为,有四种可供性(affordance)能够触发用户认知启发式(heuristic),并帮助其进行判断,包括形态、能动性、交互、导航^[14]。Hilligoss 和 Rieh 提出的可信度评估整合框架认为评估过程涉及建构层、探索层、交互层三个层次,分别对应用户对可信度的定义、启发式的思考、基于信息线索的判断^[15]。

随着人工智能检索系统在实践中运用越来越广泛,在传统的信息可信度研究基础上,关于人工智能系统的用户信任和采用意愿也已成为一个新兴的研究话题。已有研究认为内容、交互体验、情感是影响人工智能系统可信度的重要因素。在内容方面,用户感知的信息质量取决于人工智能系统向其提供的内容^[16]。用户对系统的整体信任取决于它是否能够提供准确、真实、完整和时效性高的内容来支撑用户的信息获取任务^[17]。在交互体验方面,用户通常根据人工智能系统功能特征的可靠性、灵活性、可访问性和及时性来评估系统的可信度^[18]。当人工智能系统易于访问并且出错率较低时,用户将对其给予较高的信任^[19]。在情感方面,人工智能的拟人化是由其较强的理解力和创新能力所支持的,而这样的能力使得系统能够理解与用户互动中的微妙之处^[20],进而培养用户对系统的情感信任^[21]。

纵观以上研究,可以看出,无论具体路径如何,学者对信息可信度形成的机理都持有相同的理解,即用户通过获取一定的信息线索,并将其与个人和系统交互的情境相结合,形成自身的感知和情感态度,最终对是否信任该信息做出评判。

1.2 想象可供性

可供性一词最早来源于生态心理学,关注动物与环境之间的关系,即环境能够为动物提

供何种功能^[22]。这一概念后来被引入传播学等领域,动物与环境的关系随之被引申为人与技术的关系。学者据此提出了技术可供性这一概念,指某项技术在由人们支配时,被用来达成某种效果的可能性^[23]。

学界据此进一步探讨技术为社会实践提供的可能性,这样的研究视角突破了具有显著局限性的技术决定论与社会建构论,更加关注人与技术的交互关系,强调技术不是决定性因素,最终的效能发挥取决于使用者的行为。可见,学者对于信息可信度形成机理的共识与可供性这一概念的内涵不谋而合,即都认为用户是人与技术交互过程中的重要一环^[24]。技术是为了增强人的能力,而非取代人,在迈向可信人工智能的道路上,用户是不容忽视的主体^[25]。

然而,在技术可供性的视角下,即使大多数研究关注到人与技术的关系,也仅仅是从技术的“功能性”本身出发,聚焦于技术是如何向人提供功能或进行约束的,对人与技术双向关系的观察停留在技术驱动维度,而忽略了用户感知维度^[26]。因此,Nagy 和 Neff 在 2015 年提出了“想象可供性”这一概念,指出用户往往对技术具有特定的期望,这种期望在新兴技术的使用中尤其常见,会塑造用户对技术的感知和使用行为^[27]。想象可供性聚焦人的主观能动性,在对人与技术关系的考察中,重视用户感知的作用,认为用户通过物质性、中介体验、情感来认知技术的意义。具体而言,用户将技术视作提供重要功能的实体,通过平台中介产生特定的交互体验,并且对技术产生一定的情感。

人工智能系统可信研究中对于内容、交互体验、情感三个方面的关注,又与想象可供性提出的物质性、中介体验、情感这三个用户感知技术的维度一一对应。可见,由可供性延伸而来的想象可供性与本研究拟探讨的生成式智能搜索结果可信度具有较强的相关性。将想象可供性的三个维度置于生成式智能搜索的技术研究语境中,则可作如下理解:物质性即生成式智能搜索能够为用户提供的内容搜索功能,中介体

验即用户在搜索过程中与生成式智能搜索平台的问答交互体验,情感即用户在搜索过程中被生成式智能搜索激发的情感态度。

从想象可供性的三个维度出发,结合人工智能系统可信度的已有研究^[17-18,20]及 *Nature* 对生成式智能搜索的评价^[1],本研究总结了生成式智能搜索的三个突出特征,分别为内容整合性、问答交互性、工具拟人性。其中,内容整合性对应物质性维度,指生成式智能搜索让用户能够直接获取经过加工整合的答案,而不仅仅是在搜索后收到一系列链接列表;问答交互性对应中介体验维度,生成式智能搜索可以与用户进行多回合的交互对话,支持持续性的语义检索;工具拟人性对应情感维度,与传统的搜索引擎相比,生成式智能搜索拥有拟人化的语句表达,能够与用户进行情感互动。

1.3 理论基础——SOR 模型

刺激—机体—反应 (Stimuli - Organism - Response, SOR) 模型是环境心理学领域常用的研究模型。其中,刺激(S)指来自环境的刺激,机体(O)指个体的内部感知,反应(R)指个体的行为反应^[28]。SOR 模型指出,各种类型的外部因素构成了外在环境,刺激个体产生心理感知活动,最终促使用户产生具备趋向性的行为反应。该模型往往用于理解用户行为背后的原因,在消费者行为研究领域被广泛应用。

近年来,SOR 模型的应用范围持续拓展至管理学、教育学等领域,并被引入不同的用户行为情境中。在图书情报领域,SOR 模型被用于

解决用户信息行为的相关问题,尤其适用于研究在线信息系统的用户使用行为^[29]。在应用范围不断扩大的同时,SOR 模型的内容也随之丰富,外部刺激从传统环境延伸到在线环境,机体则包含更多的感知状态,如用户的价值感知、有用性感知等,用户的行为反应也扩展到使用行为、参与行为、评论行为等^[30]。

SOR 模型基于从刺激、机体到反应的个体心理链路^[31],为厘清外部环境、个体感知及行为反应之间的关系提供了理论依据。而用户对生成式智能搜索结果的信任形成机理则与 SOR 模型的这一内在逻辑高度契合,即需要充分考虑用户作为技术使用者的角色,探讨生成式智能搜索结果的外在刺激是否会触发用户对其信任度的变化。可见,SOR 模型对于研究生成式智能搜索结果可信度,具有较强的可行性与适配性。引入 SOR 模型,在理论的指导下凝练相关要素,也能使生成式智能搜索可信度的研究过程更为规范化、体系化^[32]。

因此,本研究依据 SOR 模型,将生成式智能搜索结果界面的信息线索作为外部刺激;基于上文总结的生成式智能搜索的突出特征,将用户对内容整合性、问答交互性、工具拟人性三类特征的感知,以及与特征感知相关的具体感知属性作为机体主要内容;在反应层面,用户的内在感知将通过影响其对生成式智能搜索的信任程度,使用户产生不同的情绪或行为反应。生成式智能搜索结果可信度的 SOR 分析框架如图 1 所示。

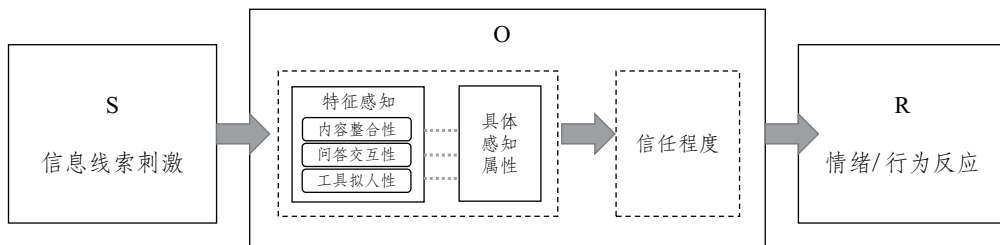


图 1 生成式智能搜索结果可信度的 SOR 分析框架

从 SOR 模型的视角出发,构建生成式智能搜索可信度的分析框架,能够较为全面地探索其形成机理,挖掘影响生成式智能搜索可信度的因素,深入描摹用户的心理或行为路径,为生成式智能搜索的优化发展提供参考。

2 研究设计

2.1 信息线索刺激标签

本研究选取生成式智能搜索的典型代表 New Bing 开展实验,探讨生成式智能搜索可信度的形成机理模型。线索效用理论内涵与可信度形成机理有着相似的路径,认为用户会根据产品内部或外部线索,形成对于产品的感知和评估。其中内部线索是产品本身固有的特征,外部线索是与产品相关的其他线索^[33]。用户对信息可信度的评判也是根据信息的内部和外部

多个线索进行的。

因此,本研究将线索效用理论引入生成式智能搜索情境中,形成 New Bing 搜索结果界面的信息线索刺激(S)标签。根据线索效用理论对内部线索和外部线索的定义,内部线索为回答内容本身相关的线索,外部线索为与搜索结果相关的其他线索。用户在多线索的情境下,会综合考量内部和外部线索,对搜索结果的可信度做出判断。

可信度线索是有助于用户进行信任判断的任何信息^[34]。结合人机交互中信息可信度的相关研究,本研究将 New Bing 搜索结果界面的常见元素分为七类信息线索刺激标签,如表 1 所示。其中,定制响应类线索、主观对话类线索、语言逻辑类线索被定义为内部线索;参考来源类线索、运行解释类线索、回答操控类线索、补充功能类线索被定义为外部线索。

表 1 信息线索刺激标签

信息线索类别		具体信息线索
内部线索	定制响应类线索	A1 回答内容与用户需求的对应性
		A2 结合上下文回答
		A3 接受调教以优化回答
	主观对话类线索	B1 情绪化的语句
		B2 承认错误的语句
		B3 带有极强主观意识的语句
		B4 在对话的开始与结尾有互动性高的问候语、结束语
	语言逻辑类线索	C1 加粗重要内容,采用符号辅助表达
		C2 回答语言包装水平
		C3 回答内容逻辑及组织条理
		C4 语句一致性
外部线索	参考来源类线索	D1 参考来源链接质量
		D2 参考来源数量及类型广度
		D3 参考来源内容与实际回答的对应性
		D4 参考来源标注与链接的对应性
		D5 参考来源内容时效性
	运行解释类线索	E1 展示响应过程
		E2 解释内容生成机制
		E3 展示检索词分解过程

续表

信息线索类别		具体信息线索
外部线索	回答操控类线索	F1 支持反馈
		F2 响应可操控
		F3 支持随时切换主题
		F4 支持调整对话模式
		F5 拒绝回答或强行结束
	补充功能类线索	G1 支持多模式回答
		G2 支持多语言
		G3 提供聊天记录,点击即可回顾问答内容
		G4 推荐用户检索内容的相关问题
		G5 推荐 New Bing 自身相关问题

这些类别的线索在已有研究中均被认为是人工智能系统信息可信度的影响因素。在对话式的人机交互中,定制响应类线索和主观对话类线索被认为是影响用户体验的主要因素,定制响应类线索是在响应生成过程中提供与用户特征和对话上下文相关的特定回答^[35],主观对话类线索是在对话过程中嵌入类似人类的语言风格、行为和信号^[36]。此外,智能系统中的语言逻辑清晰一致^[37],信息内容具备可追溯的来源数据^[38],提供关乎系统运行的解释^[39],允许人类参与系统决策并予以控制^[40],在服务层面拥有额外的便捷功能^[41]等,将增进用户对人工智能系统中信息的信任程度。

2.2 研究方法

本研究采用的方法之一为卡片分类法,即让用户将刺激标签归类至生成式智能搜索的特征感知中。卡片分类法是一种交互式研究方法,旨在揭示信息系统用户关联和分类概念的心智模型,研究参与者需要将人工制作的卡片按主题或类别进行分类^[42]。卡片分类法包括封闭式卡片分类法和开放式卡片分类法两种,前者的卡片标签及类别由研究人员预先定义,后者则由参与者自由定义。卡片分类法通常用于调查用户对网站信息分类的认知。传统的心理学研究认为,视觉卡片分类有助于构建启发性

的心理模型^[43],卡片分类法能够使研究充分考虑到参与者的观点^[44]。因此,也有学者将其拓展至其他类型的感知或行为研究中,例如,Saunders等采用卡片分类法对组织内员工的信任进行实证研究^[45],匹兹堡大学的一项研究采用卡片分类法评估学术图书馆员对研究开放性的看法^[46]。在研究群体的感知等主观问题时,卡片分类法的适用性已经得到证明。

本研究采用的方法之二为叙事研究法,基于用户讲述的生成式智能搜索使用经历,深挖其信任的形成机理。叙事研究法是指研究参与者通过讲述故事的形式表达他们关于研究主题的经验 and 观点^[47]。一方面,在叙述事件时,参与者会积极组织一个故事来讲述过去的行为,个体如何叙述其中的情节,潜在反映了影响他们感知和行为的因素^[48]。另一方面,故事本身不仅仅是对事件的重述,更能从参与者的角度向研究人员揭示其经历、感知和行为之间的关系,强调研究人员和参与者之间的关系协同^[49]。由此可见,叙事研究法尤其适用于以人为中心的日常信息行为研究,有助于从用户的角度概念化和理论化个体的信息搜寻行为。

本研究结合封闭式卡片分类法与叙事研究法开展实验。首先,每个参与者都可以使用自己的标准对固有的刺激标签卡片进行分类,可归至的类别由研究人员参照已有研究设定,包括“通常使你对搜索结果信任”“通常使你对搜

索结果不信任”“没有影响”“有影响,但影响的正负性通常不同”^[45],以反映他们对每个信息线索的信任感知。其次,在卡片分类的基础上,主要通过参与者讲述故事的方式,开展定性访谈。在实验过程中,卡片分类法与叙事研究法的结合为本研究提供了四项增益。第一,通过卡片分类有序定量数据和叙事访谈定性数据的互相补充,提升了研究的解释能力^[50];第二,卡片的视觉、触觉刺激与访谈的融合,能够引起参与者更多的使用回忆,促使其进行更深入的元认知反思;第三,在卡片分类的同时使用访谈,可确保研究人员和参与者对卡片内容的理解相同,减少歧义^[51];第四,卡片分类后接续的访谈使得实验具有更大的灵活性,允许参与者按照自己的节奏进行反思或修改,研究人员也有机会提出更多相关问题^[52]。

2.3 实验流程

本研究在正式实验之前,招募了两名具备 New Bing 使用经验的研究生参加预实验,并根据参与者在预实验结束后反馈的建议,对部分实验细节及访谈提纲进行修改。随后,共招募 31 名本科生和研究生参加正式实验,男女参与者所占比例分别为 38.71% 和 61.29%;参与者来自不同的学科,覆盖人文社会科学和自然科学;参与者的年龄在 19 至 30 岁之间;参与者均为数字原住民,具备丰富的信息检索系统使用经验,其中有 30 人使用 New Bing 的频率为每月 2—3 次及以上。本实验根据自愿参与的原则进行,并保证实验参与者的信息匿名化处理,录音在经过许可后使用。

实验开始之前,研究人员向实验参与者介绍了研究目的、实验流程及实验数据的处理方式,请参与者认真阅读并签署了知情同意书,然后向参与者演示卡片分类的操作。具体的实验流程包括以下五个主要步骤。

第一步,要求参与者填写一份前测问卷,以收集参与者的基本信息及 New Bing 的使用频率,基本信息主要包括性别、年龄、在读学位、专业等。

第二步,参与者对 29 张卡片进行分类,每张卡片代表一个信息线索,由一行解释性文字及对应的界面元素示例图组成。参与者需要自行判断这些卡片上的信息线索对搜索结果可信度是否有影响,并将这些卡片分类至研究人员设定的四个类别中。

第三步,参与者被要求回答,他们认为对搜索结果可信度有影响的信息线索,分别通过刺激内容整合性、问答交互性、工具拟人性三类特征感知中的哪一类或哪几类,进而影响了他们对搜索结果的信任。

第四步,参与者被要求回答,根据自身的搜索经历,在三类特征感知维度上,分别更信任 New Bing 还是传统搜索引擎。若有信任差异,参与者则被要求通过讲故事的方式说明是从哪些具体方面产生的。

第五步,参与者被要求对三类特征感知对搜索结果可信度的影响程度进行排序,并说明当在某一类特征感知的维度上,对搜索结果的信任程度有所降低时,通常会有怎样的反应,以及采取怎样的应对行为。

2.4 数据处理

针对卡片分类得到的有序定量数据,本研究根据用户分类结果,统计每个信息线索对搜索结果可信度有影响和对搜索结果可信度无影响两项的频次。同时,若用户认为某个信息线索通过刺激某类特征感知,进而影响了搜索结果可信度,则为该类特征感知赋值 1 分,据此分析搜索结果的各类信息线索,分别通过刺激哪类特征感知影响用户对结果的信任程度。

针对以叙事为主获得的访谈定性数据,在录音转文字并记录关键词句的基础上,采用开放编码分析访谈数据,提炼相关概念和因素,揭示用户的内在感知和行为反应。此外,针对用户认为哪类特征感知更能影响搜索结果可信度的排序数据,以及用户在各特征感知维度上更信任 New Bing 或传统搜索引擎的选择数据,进行频次统计。

3 生成式智能搜索结果可信度的 SOR 模型

3.1 刺激标签的作用分析

根据生成式智能搜索结果可信度的 SOR 分析框架,将用户对内容整合性、问答交互性、工具拟人性三类特征的感知(O1),以及与特征感知相关的具体感知属性(O2)作为机体。在机体标签提取阶段,本研究主要解决两个问题:一是三类特征感知(O1)分别受到哪些信息线索刺激(S)的影响,进而

影响搜索结果的可信度;二是三类特征感知(O1)分别与哪些具体的感知属性(O2)相关。

在探讨三类特征感知分别受到哪些信息线索刺激的影响前,需要基于卡片分类结果,确定对搜索结果可信度有影响的信息线索。在图2中,折线对应的右侧数据为本次实验中用户认为特定的信息线索对搜索结果可信度有影响的频次。尽管参与者对卡片的分类可能有所不同,但他们在感知方面的共性可以体现出来。本研究将大多数参与者归类至有影响(频次在20及以上)的刺激标签作为影响搜索结果可信度的信息线索。

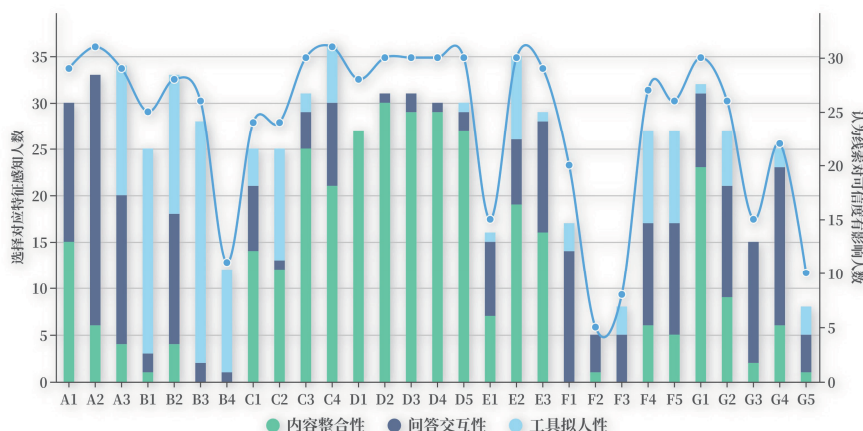


图2 信息线索刺激标签作用情况

其中,就内部线索而言,影响搜索结果可信度的信息线索包括:定制响应类线索中的回答内容与用户需求的对应性(A1)、结合上下文回答(A2)、接受调教以优化回答(A3);主观对话类线索中的情绪化的语句(B1)、承认错误的语句(B2)、带有极强主观意识的语句(B3);语言逻辑类线索中的加粗重要内容,采用符号辅助表达(C1)、回答语言包装水平(C2)、回答内容逻辑及组织条理(C3)、语句一致性(C4)。

就外部线索而言,影响搜索结果可信度的信息线索包括:参考来源类线索中的参考来源链接质量(D1)、参考来源数量及类型广度(D2)、参考来源内容与实际回答的对应性(D3)、参考来源标注与链接的对应性(D4)、参考来源内容时效性

(D5);运行解释类线索中的解释内容生成机制(E2)、展示检索词分解过程(E3);回答操控类线索中的支持反馈(F1)、支持调整对话模式(F4)、拒绝回答或强行结束(F5);补充功能类线索中的支持多模态回答(G1)、支持多语言(G2)、推荐用户检索内容的相关问题(G4)。

在将刺激标签归类至影响哪类特征感知时,大多数参与者选择了三类特征感知中的一类,图2中柱状图的三种颜色对应的左侧数据反映了选择对应特征感知的人数。在前文总结的影响搜索结果可信度的信息线索中,除了补充功能类的部分线索对应最多的特征感知存在差异外(G1对应最多的特征感知为内容整合性,G2、G4对应最多的为问答交互性),其余类别中每个信息线

索对应最多的特征感知均一致。语言逻辑类线索、参考来源类线索、运行解释类线索通过刺激用户对于内容整合性特征的感知,进而影响搜索结果的可信度;定制响应类线索、回答操控类线索通过刺激用户对于问答交互性特征的感知,进而影响搜索结果的可信度;主观对话类线索通过刺激用户对于工具拟人性特征的感知,进而影响搜索结果的可信度。

3.2 机体标签提取

本研究通过开放编码,分析以用户叙事为主的访谈数据,对31份访谈材料进行机体标签提取,剖析三类特征感知(O1)的具体感知属性(O2)。在饱和度检验方面,研究人员在访谈至

第26位参与者时未发现可提取新机体标签的内容,随后又选择第27至31位参与者进行验证,仍未出现新的要素,证明访谈材料已经足够。

通过访谈数据分析,本研究提取出与三类特征感知对应的感知属性(见表2)。用户在内容整合性维度对搜索结果的信任差异,与数据源可选择性、参考来源可靠性、结果准确性、回答透明性、结果质量判断难度、结果全面性、结果组织性相关;在问答交互性维度对搜索结果的信任差异,与检索门槛、检索持续性、语义理解能力、检索优化能力、回答稳定性相关;在工具拟人性维度对搜索结果的信任差异,与情感类人性、对话态度、思维自主性、对话专业性、安全性相关。

表2 访谈数据提取结果

特征(O1)	对应感知属性(O2)	访谈数据示例(P代表受访者)
内容整合性	数据源可选择性	“只能提供整合好的,比较局限,不能从不同的数据源选择。”(P2) “更方便了,但把我限制住了,传统搜索虽然很散,但是很全,条目更多,选择更多。”(P21)
	参考来源可靠性	“提供的知乎等链接网址没有说服力。”(P3) “来源比较权威,我会觉得靠谱,但有一次是百度知道,可信度反而降低。”(P25)
	结果准确性	“为了整合而整合,不能保证正确率。我自己的检索更准确、有保障。”(P5) “百度前面全是广告,后面开始说相关作文多少篇,很多都是无效结果,New Bing会给相关内容。”(P6)
	回答透明性	“直接展现结果,不知道为什么这么整合,不知道它背后搜索出来的全局是怎样的。”(P5) “又问它同样的问题,回答就不一样,不知道是不是时间变了,后台数据变化了。”(P22) “百度给的结果有的新、有的旧,有的是热度比较高的,摸不准,更知道New Bing是怎么推荐的。”(P24)
	结果质量判断难度	“点开参考链接看到没问题就行了,传统搜索出来的结果甄别难度很大。”(P7) “后期的查证成本太高了。”(P28)
	结果全面性	“获得的信息是它加工整合的,怀疑丢失了一些原有信息。”(P24) “整合了很多没想过的来源,一开始只想到去论文里找,New Bing拓宽了检索的信息源渠道。”(P30)
	结果组织性	“呈现形式上是结构化的,整合得很好,让人觉得有逻辑和条理。”(P21) “传统搜索只做了推荐排序,New Bing把相关的信息呈现出来,做一些梳理,可以直接用。”(P30)

续表

特征(O1)	对应感知属性(O2)	访谈数据示例(P代表受访者)
问答交互性	检索门槛	“很考验问法,有时候得到的结果不行。”(P4) “本来查全、查准需要很多的检索成本,花心思设计检索式,New Bing 不用,感觉很可靠。”(P21)
	检索持续性	“进行很多次对话,来调整搜索的结果,也深挖我的需求。”(P1) “回答具有上下文的关联性。”(P6)
	语义理解能力	“在传统搜索引擎不知道怎么问,它可以帮我拆解需求。”(P5) “传统搜索是基于词匹配,它是有针对性地按照思路回答。”(P23) “能够理解语句的细微差别,满足需求,给出完善的答案。”(P27)
	检索优化能力	“它可以不断自行修正检索式,不需要我手动修正。”(P8) “会给我允许反馈的空间,这个反馈可以帮助我下次提问。”(P11)
	回答稳定性	“回答速度慢,还经常拒绝回答。”(P16) “重复之前回答的内容,遇到这种情况挺恼火的。”(P19)
工具拟人性	情感类人性	“心里感受到对方有感情,不是冷冰冰的机器,类似于淘宝客服,像真实的人,更倾向信任它。”(P4) “情绪化的东西让我觉得很可怕,更不信任。”(P13)
	对话态度	“不是很耐心,阴阳怪气,还敷衍我。”(P7) “像是一个比较可靠的合作伙伴,在尽力帮忙找信息。”(P19)
	思维自主性	“让我觉得它是有智慧的生命体,很有创造力,很会创新。”(P3) “有人的思想,给我提供的结果有偏向性、引导性。”(P15)
	对话专业性	“并不是想获得情感价值,拟人性盖过了精确度。”(P13) “在一个很严肃的场景下,不应该用一个不太正式的回答。”(P17) “没那么客观,不是‘理中客’。”(P27)
	安全性	“传统搜索中我是操纵者,现在我被拟人化的事物所控制。”(P14) “不会问它很隐私的问题,特别担忧学术研究点的泄露。”(P6)

3.3 反应标签提取

搜索结果界面的信息线索通过刺激与多项感知属性相关的特征感知,进而影响用户对搜索结果的信任程度,而三类特征感知维度下的可信度降低会导致用户产生不同的反应(R)。本研究基于对访谈材料的分析,提取出该阶段的用户反应标签。

在内容整合性特征感知的维度上,当用户对搜索结果的信任程度有所降低时,通常产生失望性质的情绪,在行为上,往往选择结束当前对话,转而使用传统搜索引擎,或将传统搜索引擎和生成式智能搜索结合起来使用。

在问答交互性特征感知的维度上,当用户对搜索结果的信任程度有所降低时,通常在情绪方面选择理解搜索工具,在行为上以重新构

造检索式、继续调整搜索工具为主。

在工具拟人性特征感知的维度上,当用户对搜索结果的信任程度有所降低时,在情绪上具有两极分化的情况,一种是感到愤怒,另一种则是感到有趣。在前者的情绪主导下,用户通常选择结束使用搜索工具,小部分用户会在对话中进行语言还击;在后者的情绪主导下,用户则会试图继续对话并观察事态发展,后续以谈资的形式在线下分享该经历。

3.4 SOR 模型构建

本研究依据前期形成的生成式智能搜索结果可信度 SOR 分析框架,基于提取归纳的信息线索刺激标签(S)、特征感知标签与具体感知属

性标签(O)、情绪及行为反应标签(R),构建了最终的SOR模型,如图3所示。

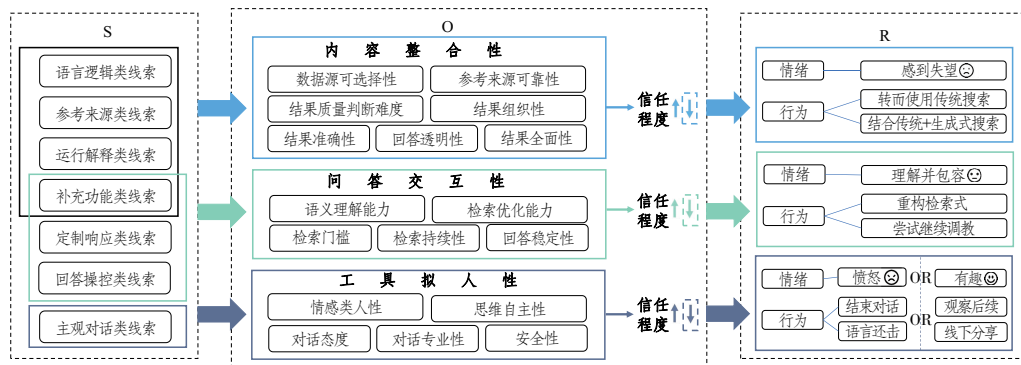


图3 生成式智能搜索结果可信度的SOR模型

在信息线索刺激用户感知的环节,各类别线索与其刺激特征感知的具体感知属性息息相关。如刺激内容整合性特征感知的语言逻辑类线索、参考来源类线索、运行解释类线索分别对应内容整合性相关的结果组织性、参考来源可靠性和回答透明性;刺激问答交互性特征感知的定制响应类线索和回答操控类线索分别对应问答交互性相关的检索优化能力和回答稳定性;刺激工具拟人性特征感知的主观对话类线索对应工具拟人性相关的感知属性。此外,在具有刺激性的线索中,外部线索更多是通过刺激内容整合性发挥作用,这表明用户将非检索内容的元素与内容本身相结合,以综合评价搜索结果。

在特征感知影响信任程度的环节,根据用户对三类特征感知影响程度的排序结果(见图4左),大多数用户认为内容整合性是最影响可信度的(排序为1的次数最多),问答交互性其次(排序为2的次数最多),工具拟人性对可信度的影响最小(排序为3的次数最多)。

但在实际的使用体验中,根据用户在三类特征感知维度上更信任生成式智能搜索还是传统搜索引擎的选择结果(见图4右),可知生成式智能搜索在问答交互性维度上让用户最为信任,在对可信度影响最大的内容整合性维度上,生成式智能搜索获取的用户信任反而并不突出。可以看出,用户对重视的特征感知的排序,与实践中的可信度高低排序并不相符。

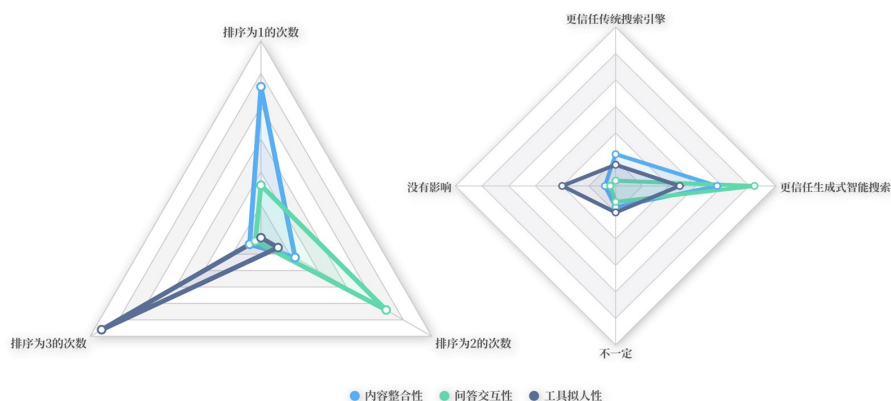


图4 各特征感知对信任的影响程度、对应维度的实际可信度

在信任程度影响个体反应的环节,三类特征感知的具体变化通过影响用户对搜索结果的信任程度,引发了不同的情绪和行为反应。用户认为内容整合性最影响可信度,相应地,内容整合性维度上可信度的降低导致的用户反应较为负面。针对问答交互性维度上的可信度降低,用户通常予以包容性反应,尝试调整个人的检索策略以优化结果。工具拟人性维度导致的两极分化反应也与用户对该特征的感知息息相关。由图4可知,大部分用户在工具拟人性方面的信任倾向分为两类:一类是更信任生成式智能搜索,这类用户往往会对工具拟人性维度上的可信度降低感到愤怒,情绪反应较为激烈;另一类则认为工具拟人性对于搜索结果可信度没有影响,这类用户通常将生成式智能搜索拟人性方面的瑕疵作为趣事看待,选择忽视其可能对搜索结果可采纳性造成的负面影响。

4 启示与展望

生成式智能搜索是提供信息服务的新兴工具,研发者在对其的优化与更新中需要把握用户对搜索结果可信度的评估机制,根据用户检索体验催生的实践需求来完善系统。本研究以SOR模型为基础,聚焦生成式智能搜索的三类特征感知,从影响它们的信息线索、具体感知属性、各特征感知对可信度的影响、可能导致的情绪或行为反应等方面进行挖掘与探讨,总结出优化生成式智能搜索的几点方向。

第一,确保获得的内容真实有效是检索过程中优先级最高的需求。生成式智能搜索的内容整合性节省了用户初始检索花费的时间,但在查证方面损耗了用户更多的精力。相比于单一的聊天机器人,大多数认为生成式智能搜索更可信的用户的出发点在于其具备参考来源,可以直接在检索结果界面点击链接查证。因此,参考来源是生成式智能搜索结果可信最重要的保障元素,研发者可以为参考来源设置限制条件,保障来源与搜索结果一一对应,避免出

现只有回答而没有出处。要保障底层模型优先抽取时效性更高的支撑数据,允许登录用户设置个性化标签,选择倾向采用的数据源类型,如不使用知乎、百度知道数据等。

此外,在查证过程中,用户常通过传统网页搜索来检验生成式智能搜索结果可信度,研发者不用回避这种做法。问答与网页搜索相结合本就是生成式智能搜索区别于聊天机器人的特点之一,研发者应致力于优化用户将两者结合的使用体验。当前许多用户不常在生成式智能搜索上通过问答界面与网页搜索界面切换的方式进行查证,原因即在于实践中存在问答缓存丢失的问题,用户从问答界面切换至网页搜索查证后,再度返回问答界面时,系统已经自动开启新对话,因此,应注重改善网页搜索和问答搜索切换的便捷性。同时,除了PC端分屏查证外,多设备查证的方式也受到用户青睐,而当前生成式智能搜索移动设备端的使用便捷性较低,研发者应注重生成式智能搜索移动设备端的优化。

第二,生成式智能搜索在问答交互性方面获得较高评价,但其搜索结果界面的运行解释类元素被用户认为通过刺激内容整合性影响信任,表明其解释仅被用户作为额外的补充信息,并未从根本上提升用户的交互体验。已有研究认为,可解释性由清晰的、具有解释作用的细节信息所保障^[53]。而生成式智能搜索提供的解释信息作用较小,主要是由于其他方面的不匹配性,如不同时间对同一个问题的回答不一致,回答操控类元素缺乏反馈,回答不稳定等,都会降低用户对可解释性的评价。

从生成式智能搜索的特殊性出发,可解释性除了与表面的解释信息相关外,还蕴含在回答的潜在表现中,用户感知到的结果客观性、统一性对于可解释性的提升具有补充作用。因此,研发者需要优化回答的稳定性,保证针对同一问题抓取的底层数据一致,增加回答操控类元素的反馈,降低在缺乏明确理由下系统自行结束对话的概率。这一点也可以拓展至其他人

工智能系统的可解释性研究中,除注重解释性信息的融入外,也应关注其运行机制一致性和稳定性的用户感知。此外,构造生成式智能搜索检索式对用户而言是一大难题,用户认为系统能够自行优化检索式,但这种优化是隐性的,仅偶尔体现在运行解释类信息中,研发者若在实践中将检索式的优化方向显性化,如向用户推荐调整的检索语句,则更能获取用户的信任。

第三,已有研究认为拟人性会增强对话式人工智能的可信度,用户倾向于认为具备人类表达风格的智能系统更值得亲近和信任^[54]。但本研究发现,检索过程中拟人性融入的不恰当,会导致用户对生成式智能搜索结果的可信度降低。这与生成式智能搜索的应用情境高度相关,此前的对话式人工智能并未深入应用于信息检索情境中,而本实验中,大多数用户使用生成式智能搜索的主要目的在于获取精确相关的信息,而非单纯的娱乐消遣。一方面,在获取专业性信息的

情境下,拟人化的表达将使用户降低对系统专业性的评价;另一方面,当系统在对话中返回情绪化或态度不佳的语句时,会降低用户对其提供信息的信任度。因此,研发者应在非情感类问题的检索情境下,调整拟人性相关参数,注重把控系统拟人化程度,避免不合时宜的类人表达。

总而言之,每类特征感知对于用户的影响是不同的,在同一特征维度上,用户对生成式智能搜索可信度的评价不同,而在不同维度上,同一用户的信任程度也存在差异。这与用户的个人背景、认知习惯相关,也与具体的使用情境、信息需求密不可分。但在本质上,用户对于搜索结果的核心诉求是相同的,可靠、专业、精准的内容才能获得用户青睐。在生成式智能搜索当下的发展阶段中,唯有聚焦用户使用中的检索需求,才能让搜索结果的可信度得到质的提升,真正实现个性化智能搜索,让信息检索在变革性技术的加持下走向更加高效便捷的未来。

致谢:本文系国家自然科学基金创新群体项目“信息资源管理”(项目编号:71921002)和湖北省自然科学基金创新群体项目“以人为本的人工智能创新应用”(项目编号:ZRQT2023000026)的研究成果。

参考文献

- [1] STOKEL-WALKER C. AI chatbots are coming to search engines—can you trust the results?[J/OL]. *Nature*, 2023[2023-07-18]. <https://www.nature.com/articles/d41586-023-00423-4>.
- [2] CASTELO N, BOS M W, LEHMANN D R. Task-dependent algorithm aversion[J]. *Journal of Marketing Research*, 2019, 56(5): 809-825.
- [3] KOCON J, CICHECKI I, KASZYCA O, et al. ChatGPT: jack of all trades, master of none[J/OL]. *Information Fusion*, 2023[2023-07-18]. <https://arxiv.org/pdf/2302.10724.pdf>.
- [4] DWIVEDI Y K, KSHETRI N, HUGHES L, et al. “So what if ChatGPT wrote it?” Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy[J/OL]. *International Journal of Information Management*, 2023, 71[2023-07-18]. <https://www.sciencedirect.com/science/article/pii/S0268401223000233>.
- [5] 于水, 范德志. 新一代人工智能(ChatGPT)的主要特征、社会风险及其治理路径[J]. *大连理工大学学报(社会科学版)*, 2023, 44(5): 28-34. (YU S, FAN D Z. The main characteristics, social risks and governance paths of the new generation of artificial intelligence(ChatGPT)[J]. *Journal of Dalian University of Technology (Social Sciences)*, 2023, 44(5): 28-34.)
- [6] HACKER P, ENGEL A, MAUER M. Regulating ChatGPT and other large generative AI models[C]//*Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency*. Chicago, USA, 2023: 1112-1123.
- [7] 宋小康, 赵宇翔, 宋士杰, 等. 社会技术系统范式下 AI 赋能的替代信息搜索: 特征、理论框架与研究展望[J]. *图书情报知识*, 2023, 40(4): 111-121. (SONG X K, ZHAO Y X, SONG S J, et al. The features, theoretic-

- cal framework and research prospects of AI-enabled surrogate information searching;a sociotechnical system paradigm[J]. *Documentation,Information & Knowledge*,2023,40(4):111-121.)
- [8] 吴丹,唐源. 搜索引擎结果页面(SERP)研究述评[J]. *情报学报*,2018,37(2):220-230. (WU D,TANG Y. A review on the Search Engine Result Page(SERP)[J]. *Journal of the China Society for Scientific and Technical Information*,2018,37(2):220-230.)
- [9] 宋士杰,赵宇翔,朱庆华. iField 视域下的信息可信度研究:概念溯源、主题演化与未来展望[J]. *中国图书馆学报*,2022,48(1):107-126. (SONG S J,ZHAO Y X,ZHU Q H. Information credibility research in the iField:conceptual development,topic evolution,and future direction[J]. *Journal of Library Science in China*,2022,48(1):107-126.)
- [10] 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. (WANG P,CHENG Q K. Review and progress in research on credibility evaluation of information on the Web [J]. *Journal of Information Resource Management*,2013,3(1):46-52.)
- [11] PIROLI P,CARD S. Information foraging[J]. *Psychological Review*,1999,106(4):643-675.
- [12] FOGG B J. Prominence-interpretation theory:explaining how people assess credibility online[C]//CHI'03 Extended Abstracts on Human Factors in Computing Systems. New York,USA:Association for Computing Machinery,2003:722-723.
- [13] FLANAGIN A J,METZGER M J. The role of site features,user attributes,and information verification behaviors on the perceived credibility of Web-based information[J]. *New Media & Society*,2007,9(2):319-342.
- [14] SUNDAR S S. The MAIN model:a heuristic approach to understanding technology effects on credibility[M]//METZGER M J,FLANAGIN A. *Digital media,youth,and credibility*. Cambridge:The MIT Press,2008:73-100.
- [15] HILLIGOSS B,RIEH S Y. Developing a unifying framework of credibility assessment:construct,heuristics,and interaction in context[J]. *Information Processing & Management*,2008,44(4):1467-1484.
- [16] ASHFAQ M,YUN J,YU S,et al. I,Chatbot:modeling the determinants of users' satisfaction and continuance intention of AI-powered service agents[J/OL]. *Telematics and Informatics*,2020,54[2023-07-18]. <https://www.sciencedirect.com/science/article/abs/pii/S0736585320301325>.
- [17] KIM J,GIROUX M,LEE J C. When do you trust AI? The effect of number presentation detail on consumer trust and acceptance of AI recommendations[J]. *Psychology & Marketing*,2021,38(7):1140-1155.
- [18] BEDUÉ P,FRITZSCHE A. Can we trust AI? An empirical investigation of trust requirements and guide to successful AI adoption[J]. *Journal of Enterprise Information Management*,2022,35(2):530-549.
- [19] TRIBERTI S,DUROSINI I,CURIGLIANO G,et al. Is explanation a marketing problem? The quest for trust in artificial intelligence and two conflicting solutions[J]. *Public Health Genomics*,2020,23(1-2):2-5.
- [20] PELAU C,DABIJA D C,ENE I. What makes an AI 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/OL]. *Computers in Human Behavior*,2021,122[2023-07-18]. <https://www.sciencedirect.com/science/article/abs/pii/S0747563221001783>.
- [21] SHI S,GONG Y,GURSOY D. Antecedents of trust and adoption intention toward artificially intelligent recommendation systems in travel planning;a heuristic-systematic model[J]. *Journal of Travel Research*,2021,60(8):1714-1734.
- [22] GIBSON J J. *The perception of the visual world*[M]. Boston:Houghton Mifflin,1950:35-38.
- [23] GAVER W W. Technology affordances[C]//CHI'91:Proceedings of the SIGCHI Conference on Human Factors in Computing Systems. New York,USA:Association for Computing Machinery,1991:79-84.
- [24] 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. (SONG S J,ZHAO Y X,ZHU Q H. From ELIZA to ChatGPT:AI-Genera-

- ted Content(AIGC)credibility evaluation in human-intelligent interactions experience[J]. Information and Documentation Services,2023,44(4):35-42.)
- [25] 吴丹,孙国烨. 迈向可解释的交互式人工智能:动因、途径及研究趋势[J]. 武汉大学学报(哲学社会科学版),2021,74(5):16-28. (WU D,SUN G Y. Towards explainable interactive artificial intelligence:motivations, approaches,and research trend[J]. Wuhan University Journal (Philosophy & Social Science),2021,74(5): 16-28.)
- [26] 杨奇光. 技术可供性视域下的数字新闻实践及其边界重思[J]. 西北师大学报(社会科学版),2023,60(4):83-91. (YANG Q G. Rethinking of digital journalism from technology affordance perspective and its boundary[J]. Journal of Northwest Normal University(Social Sciences),2023,60(4):83-91.)
- [27] NAGY P, NEFF G. Imagined affordance;reconstructing a keyword for communication theory[J/OL]. Social Media + Society,2015,1(2)[2023-07-18]. <https://journals.sagepub.com/doi/full/10.1177/2056305115603385>.
- [28] MEHRABIAN A,RUSSELL J A. An approach to environmental psychology[M]. Cambridge:The MIT Press, 1974:132-135.
- [29] ZHAN Y,LI P,QU Z,et al. A learning-based incentive mechanism for federated learning[J]. IEEE Internet of Things Journal,2020,7(7):6360-6368.
- [30] 李琪,李欣,魏修建. 整合SOR和承诺信任理论的消费者社区团购研究[J]. 西安交通大学学报(社会科学版),2020,40(2):25-35. (LI Q,LI X,WEI X J. Study on consumers' community group purchase based on the integration of SOR and commitment-trust theory[J]. Journal of Xi'an Jiaotong University(Social Sciences),2020, 40(2):25-35.)
- [31] KINI R B,BOLAR K,ROFIN T M,et al. Acceptance of location-based advertising by young consumers:a Stimulus-Organism-Response(S-O-R) model perspective[J/OL]. Information Systems Management,2023[2023-07-18]. <https://www.tandfonline.com/doi/full/10.1080/10580530.2023.2214843>.
- [32] 徐孝娟,赵宇翔,史如菊,等. SOR理论在国内图书情报学领域的采纳:溯源、应用及未来展望[J]. 情报资料工作,2022,43(5):98-105. (XU X J,ZHAO Y X,SHI R J,et al. The adoption of SOR theory in the field of library and information science in China:traceability,application and future prospects[J]. Information and Documentation Services,2022,43(5):98-105.)
- [33] 刘子萌,袁勤俭. 五大线索理论的发展及其应用进展[J]. 现代情报,2021,41(10):140-149. (LIU Z M, YUAN Q J. The development and application of five cue theories[J]. Journal of Modern Information,2021,41(10):140-149.)
- [34] LIAO Q V,SUNDAR S S. Designing for responsible trust in AI systems;a communication perspective[C]//Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency. Seoul, Korea,2022: 1257-1268.
- [35] MOUSSAWI S,KOUFARIS M,BENBUNAN-FICH R. How perceptions of intelligence and anthropomorphism affect adoption of personal intelligent agents[J]. Electronic Markets,2021,31:343-364.
- [36] GO E,SUNDAR S S. Humanizing chatbots:the effects of visual,identity and conversational cues on humanness perceptions[J]. Computers in Human Behavior,2019,97:304-316.
- [37] WOJTON H M,PORTER D T,LANE S,et al. Initial validation of the Trust of Automated Systems Test(TOAST) [J]. The Journal of Social Psychology,2020,160(6):735-750.
- [38] AHMADI V,TUTSCHKU K. Privacy and trust in cloud-based marketplaces for AI and data resources[C]//11th IFIP International Conference on Trust Management. Gothenburg,Sweden,2017:223-225.
- [39] SHIN D. The effects of explainability and causability on perception,trust,and acceptance:implications for explainable AI[J/OL]. International Journal of Human-Computer Studies,2021,146[2023-07-18]. <https://>

- www. sciencedirect. com/science/article/abs/pii/S1071581920301531.
- [40] SMUHA N A. The EU approach to ethics guidelines for trustworthy artificial intelligence[J]. Computer Law Review International,2019,20(4):97-106.
- [41] AMEEN N,TARHINI A,REPPPEL A,et al. Customer experiences in the age of artificial intelligence[J/OL]. Computers in Human Behavior,2021,114[2023-07-18]. <https://www. sciencedirect. com/science/article/pii/S0747563220302983>.
- [42] SCHMETTOW M,SOMMER J. Linking card sorting to browsing performance—are congruent municipal websites more efficient to use?[J]. Behaviour & Information Technology,2016,35(6):452-470.
- [43] FIORE S M,CUEVAS H M,OSER R L. A picture is worth a thousand connections;the facilitative effects of diagrams on mental model development and task performance[J]. Computers in Human Behavior,2003,19(2):185-199.
- [44] MCTAVISH J. Everyday life classification practices and technologies:applying domain-analysis to lay understandings of food,health,and eating[J]. Journal of Documentation,2015,71(5):957-975.
- [45] SAUNDERS M N K,DIETZ G,THORNHILL A. Trust and distrust;polar opposites,or independent but co-existing?[J]. Human Relations,2014,67(6):639-665.
- [46] LYON L,MATTEEN E,JENG W,et al. Investigating perceptions and support for transparency and openness in research;using card sorting in a pilot study with academic librarians[J]. Proceedings of the Association for Information Science and Technology,2016,53(1):1-5.
- [47] BATES J A. Use of narrative interviewing in everyday information behavior research[J]. Library & Information Science Research,2004,26(1):15-28.
- [48] MERAZ R L,OSTEEN K,MCGEE J. Applying multiple methods of systematic evaluation in narrative analysis for greater validity and deeper meaning[J/OL]. International Journal of Qualitative Methods,2019,18[2023-07-18]. <https://journals. sagepub. com/doi/10.1177/1609406919892472>.
- [49] HAYDON G,BROWNE G,VAN DER RIET P. Narrative inquiry as a research methodology exploring person centred care in nursing[J]. Collegian,2018,25(1):125-129.
- [50] EDMONDSON A C,MCMANUS S E. Methodological fit in management field research[J]. Academy of Management Review,2007,32(4):1246-1264.
- [51] SAUNDERS M N K,THORNHILL A. Researching sensitively without sensitizing;using a card sort in a concurrent mixed methods design to research trust and distrust[J]. International Journal of Multiple Research Approaches,2011,5(3):334-350.
- [52] CONRAD L Y,TUCKER V M. Making it tangible;hybrid card sorting within qualitative interviews[J]. Journal of Documentation,2019,75(2):397-416.
- [53] 孔祥维,唐鑫泽,王子明. 人工智能决策可解释性的研究综述[J]. 系统工程理论与实践,2021,41(2):524-536. (KONG X W,TANG X Z,WANG Z M. A survey of explainable artificial intelligence decision[J]. Systems Engineering—Theory & Practice,2021,41(2):524-536.)
- [54] LU L,MCDONALD C,KELLEHER T,et al. Measuring consumer-perceived humanness of online organizational agents[J/OL]. Computers in Human Behavior,2022,128[2023-07-18]. <https://dl. acm. org/doi/abs/10.1016/j. chb.2021.107092>.

吴 丹 武汉大学信息管理学院、武汉大学人机交互与用户行为研究中心教授,博士生导师。湖北武汉 430072。

孙国烨 武汉大学信息管理学院博士研究生。湖北 武汉 430072。

(收稿日期:2023-07-24)