

人工智能在虚拟化社交平台中的应用困境与化解路径

司 蕾

贵州大学哲学学院, 贵州 贵阳

收稿日期: 2025年6月10日; 录用日期: 2025年6月23日; 发布日期: 2025年7月23日

摘 要

人工智能越来越多地应用于公众所熟知的领域, 在虚拟化社交平台中, 人工智能的身影无处不在, 用户身份识别与管理、虚拟形象与场景构建、社交互动辅助、个性化内容推荐与服务是其主要应用形式。然而, 人工智能的广泛应用在维系人际关系和推动社会进步的同时, 其潜在的风险和问题也逐渐显现。通过总结和梳理人工智能在虚拟化社交平台应用中存在的隐私泄露风险与虚假信息传播、算法偏见与不公平性、社交关系异化、责任主体冲击与边界模糊等问题, 力图揭示当前虚拟化社交平台中人工智能应用的困境, 探索人工智能在未来健康发展的道路, 推动人工智能进一步向纵深发展。

关键词

人工智能, 虚拟, 社交平台

The Dilemma of Artificial Intelligence Application in Virtualized Social Platforms and the Path to Resolve It

Lei Si

School of Philosophy, Guizhou University, Guiyang Guizhou

Received: Jun. 10th, 2025; accepted: Jun. 23rd, 2025; published: Jul. 23rd, 2025

Abstract

Artificial intelligence is increasingly being applied in areas known to the public, and is ubiquitous in virtualized social platforms, with user identification and management, virtual image and scene construction, social interaction assistance, and personalized content recommendation and services

being its main forms of application. However, while the extensive application of AI maintains interpersonal relationships and promotes social progress, its potential risks and problems are gradually emerging. By summarizing and sorting out the risks of privacy leakage and false information dissemination, algorithmic bias and unfairness, alienation of social relationships, impact of responsible subjects and boundary blurring in the application of AI in virtualized social platforms, we seek to reveal the dilemmas of the current application of AI in virtualized social platforms, to explore the path of AI's healthy development in the future, and to promote AI's further development in the vertical and deeper aspects.

Keywords

Artificial Intelligence, Virtual, Social Platforms

Copyright © 2025 by author(s) and Hans Publishers Inc.

This work is licensed under the Creative Commons Attribution International License (CC BY 4.0).

<http://creativecommons.org/licenses/by/4.0/>



Open Access

1. 引言

迄今为止,人类对人工智能的应用已经越来越广泛,社会智能化程度日益加深,如扫地机器人、自动驾驶、社交平台等日常应用中都包含着人工智能技术的东西。在虚拟化社交平台上,人工智能技术的应用主要表现在虚拟身份、形象以及场景等的构建。当前,虚拟化社交平台发展态势迅猛,市场规模日益扩大,用户群体不断增长。随着人工智能技术在虚拟化社交平台当中实现更广泛、更深层的应用,用户体验和娱乐生活等都获得了极大的优化,然而,这也带来了一定的伦理失范现象,对网络安全和现实社会造成不可小觑的负面影响。

2. 人工智能在虚拟化社交平台中的应用现状

人工智能在今天作为一个与人类息息相关的概念,对人们来说并不陌生,广泛渗透在人们的生产生活中。人工智能在虚拟化社交平台中具有独特的应用形式,在现代社会的人际交往中有着重要的作用。

2.1. 人工智能和虚拟化社交平台的概述

人工智能(Artificial intelligence)是通过计算机科学与技术等多学科理论交叉发展而来的具有模拟、延伸和扩展人类智能等功能的理论、方法、技术或应用系统,基于大数据、算法与算力能为人类提供产品、服务、应用的智能体[1]。人工智能本身有着极其复杂的流变性,但在其丰富多变的表现形态之下,又蕴含着共通之处,有着共同的本质特征,包括类人性、交叉性、生成性、智能性等共性,大多数人在此基础上将其看作是生产出来以满足人类的主体需要的一种“机器”。人工智能已从理论走向广泛应用,深刻影响着医疗、交通、金融等领域的变革与发展,当前人工智能的应用以“弱人工智能”[2](Narrow AI)为主。“人工智能伦理”是对传统社会伦理的突破和超越,主要表现在人工智能伦理超出了人与人之间关系的部分,而增加了人与人工智能或技术之间的新型伦理关系,将传统伦理处理人与人的关系转变为处理以技术为中介的人与其他人、人与自身、人与社会、人与自然等的关系[3]。

虚拟化社交平台是传统人际社交延伸出来的一种数字化社交,其建立在一些先进技术手段如人工智能(AI)、虚拟现实(VR)、增强现实(AR)的基础之上,通过打造虚拟空间或场景为用户提供社交便利,能够打破距离、性别、年龄等现实因素的限制,用户能以虚拟化身形式展开多样化社交,在平台以语音和

文字进行聊天,参与线上娱乐活动(如虚拟演唱会、电竞比赛)等,对于提升用户社交体验具有重要意义。虚拟化社交平台种类繁多,目前,官方统计报告中少有对社交平台进行统计排名的数据,根据《2022 主流社交媒体平台趋势洞察报告》[4]等商业机构报告,可以将虚拟化社交平台按社交场景与功能划分为:综合社交平台,其特点是覆盖面广且服务多类群体,具有通信宣传、分享生活内容、提供日常服务等多方面的功能,例如微信;兴趣社交平台,基于共同兴趣爱好或行为习惯等建立起来的,主要服务于用户娱乐,例如抖音、哔哩哔哩;匿名社交平台,顾名思义,用户可以隐藏或重构自己的身份,以其构建的虚拟形象参与社交;商务社交平台,主要是为职场人士提供工作机会,能够彼此交流行业信息,满足其商务合作需求等。

2.2. 人工智能在虚拟化社交平台中的应用形式与积极效应

2.2.1. 用户身份识别与管理

人工智能正深刻影响着虚拟化社交平台的运行逻辑。一方面,人工智能可以快速准确地识别社交平台中用户的生物特征,进行用户身份认证,为社交平台管理用户提供坚实可靠的身份验证体系。在一些社交平台的登录和注册环节,平台会利用由人工智能构建的生物识别系统来综合采集用户的脸部特征和声音等生物信息,并将其与用户账号进行绑定,用以完成毫秒级的身份核验。另一方面,社交平台也可以通过人工智能来分析用户的行为模式,对用户身份进行管理,人工智能可以实时和持续监测用户的操作行为,包括打字快慢、点击频率、浏览历史和社交习惯等数据,进而构建起关于每一用户的行为模式模型,一旦检测到用户发生异常行为,将触发预警机制并要求用户进行额外的身份验证,这能够有效识别出被盗用的账号或虚假账号。

2.2.2. 虚拟形象与场景构建

人工智能通过构建虚拟形象和社交场景,在传统人际社交的基础上拓展了社交的形式和边界。一些短视频社交平台如抖音和快手,就利用人工智能深度学习算法来快速生成 3D 虚拟形象,用户可以通过上传自己的照片来构建专属的虚拟形象。用户也能够输入风格关键词,如“古风”“甜美”等来生成相应的表情、妆容和服饰。在直播类的社交场景中,如 B 站等平台上的虚拟主播,可以采取人工智能驱动的动作映射和表情模拟技术将主播的真实动作和表情实时转变为虚拟的动态形象。一些虚拟偶像如洛天依通过虚拟演唱会与人实现互动。虚拟形象的构建具有高度自由度,不仅能自定义身份、工作、性别、性格、年龄等,也可以重构自己的身份信息,这有助于减少在社交中因这些因素带来的隐形歧视。人工智能在构建虚拟场景上,可以根据社交主题和用户需求生成与之相符的虚拟社交空间。例如, Soul 这一社交平台可以生成咖啡馆等生活场景,QQ 和微信用户与好友进行线上视频会议时可选择将背景替换为人工智能生成的虚拟场景。这表明用户足不出户就能实现场景化社交,提升社交参与的便利性。

2.2.3. 社交互动辅助

用户在社交平台中可以使用人工智能提供的自动回复、智能提醒、语言翻译、社交陪伴等功能提升社交效率和体验。在一些用户互动频次高的社交平台如微博和抖音,有虚拟的客服机器人自动回复为用户答疑,人工智能基于自然语言处理技术可以做到实时提醒用户发送的私信和评论。在一些办公社交平台如钉钉和飞书,可以利用人工智能生成用户的会议安排、日程提醒和任务截止提醒等。微信朋友圈也利用人工智能提醒用户去点赞和评论好友的生日等重要时刻,能够避免因疏忽而遗漏与好友的重要互动。在 TikTok 和 Facebook 等全球性社交平台中,人工智能可以给不同国度的用户实时翻译语言,让用户以母语浏览他人的动态和评论,真正实现“跨国界社交”,有利于促进多元文化交流。恋与制作人和光与夜之恋等恋爱社交平台能够通过人工智能打造虚拟角色为用户提供社交陪伴,这在一定程度上可以满足用户的情感需求,提升社交互动的“温度”。另外,依托人工智能为残障人士、偏远地区人群提供平等社

交机会也是重要内容。

2.2.4. 个性化内容推荐与服务

人工智能通过算法记录和分析用户的兴趣、浏览历史和行为以及社交关系等数据，可以为用户推荐相应的视频、文章和广告等个性化内容。在微博和小红书这类社交平台，人工智能根据自然语言处理技术可以记录并分析用户发布的动态、评论以及收藏内容中的关键词和情感信息，并结合用户浏览类型、停留时长和点击频率等数据来给用户“贴标签”。如果用户对某一话题或内容进行频繁点赞，那这类标签就会被赋予该用户。用户的社交关系网络也是人工智能推荐算法的参考维度之一，通过分析用户的关注列表、参与群组以及互动频繁的好友，来向其推荐相关内容，例如，用户的大部分好友都关注了某一话题，人工智能会认为用户对这一领域存在潜在兴趣。在抖音和淘宝，人工智能结合用户的购物历史、搜索记录和社交行为等数据，可以构建用户的消费偏好模型。个性化内容推荐与服务能够极大地减少用户筛选信息的时间成本，使其更易找到兴趣相投的内容与人群。

3. 人工智能在虚拟化社交平台中应用的困境

人工智能在虚拟化社交平台中的应用确实为人们提供了便利的社交方式，具有重要的积极影响，这是不可否认的事实。但值得深思的是，人工智能在为人的虚拟化社交提供服务的过程中，存在一定的困境，集中体现在以下四个方面。

3.1. 隐私泄露风险与虚假信息传播

虚拟化社交平台中人工智能通过用户注册、日常交互和行为习惯等多环节采集数据，涵盖用户基础身份信息、消费习惯、社交关系网络等各个隐私方面，一些社交平台打着“优化服务和体验”的旗号，过度索取用户权限，甚至将用户的敏感健康信息等纳入采集范畴。用户数据可能被用于未经用户授权的商业开发，甚至被非法共享给第三方机构，导致用户隐私暴露。社交平台中海量数据多以集中式存储于云端，一旦平台存在安全漏洞，不法分子便可利用人工智能自动化攻击手段，如恶意爬虫和智能漏洞扫描，突破防护系统，窃取用户隐私。人工智能在虚拟化社交平台中的应用还存在虚假信息传播的现象。人工智能依托深度学习与自然语言处理技术，能模仿人类的创作风格，快速生成文本、图像甚至视频。不法分子可能利用这一特性炮制虚假新闻，编造名人轶事，制作夸张的虚假广告，这些内容在语言逻辑和视觉效果上可以以假乱真，普通用户极难辨别。个性化推荐算法则为这些虚假信息的传播提供了“温床”，容易导致用户沉浸在由算法构建的“信息茧房”中，难以接触到多元观点和真实信息。

3.2. 算法偏见与不公平性

虚拟化社交平台中的人工智能算法训练数据偏差产生的偏见问题日益受到重视。在虚拟化社交平台中，人工智能算法往往基于历史用户的行为与互动记录，如果这些数据本身存在性别、种族或地域等方面的偏见，就会形成恶性循环。例如，在一些算法训练数据中，女性的相关职业通常被描述为“局限于传统岗位”，这就导致算法在内容推荐时可能会减少女性创新领域就职内容的推送，进一步形成系统性歧视。同时，算法决策缺乏透明度与可解释性是另一重要问题。算法黑箱问题使社交平台用户难以理解推荐内容、账号评级等决策背后的逻辑。当算法错误封禁用户账号或因不透明的推荐机制导致用户流量骤减时，用户既无法知晓具体判定依据，也难以获得有效申诉渠道。这种不透明性不仅损害用户权益，更破坏了用户对平台的信任，长此以往，将对社交平台的健康发展产生一定的影响。

3.3. 社交关系异化

人们过度依赖人工智能辅助社交容易削弱人类的真实社交能力。当人们频繁借助聊天机器人代替与

真人的日常交流,使用人工智能生成的话术维系社交关系时,会导致人们面对面交流时观察微表情、感知语气与临场反应的能力逐渐钝化。如一些职场新人过度使用人工智能生成社交话术,在实际商务交流中,却因无法灵活应对突发状况而错失机会。这种现象印证了心理学中的“用进废退”理论,长期依赖虚拟社交工具,看似便捷的沟通方式实则是“温水煮青蛙”。人工智能在虚拟社交平台中的应用还面临情感真实性与信任关系构建的严峻挑战。在一些恋爱模拟社交平台中,人工智能凭借强大的语言、画面和声音生成能力,能模拟出细腻的情感表达和恋爱场景,让用户在虚拟世界中得到情绪共鸣。但算法根据数据模型生成的“标准化回应”是不排他的,可以满足多个人的情感需求。这恰恰违背了人类的恋爱伦理观念,因为人类的恋爱是所有人类情感中排他属性最为强烈的,很大可能会改变人类对于爱情的定义和信任程度。青少年作为虚拟化社交平台中的主力军,也是受影响最深的群体,最容易在人工智能构建的虚拟社交世界中流连忘返。在虚拟社交的“滤镜”下,部分青少年易产生容貌焦虑和社交自卑,进而引发抑郁等心理问题。并且,长时间使用电子设备参与虚拟社交,还会导致青少年视力下降和颈椎劳损等身体问题,阻碍其正常生长发育。

3.4. 责任主体冲击与边界模糊

人工智能对责任主体的冲击体现为主体道德责任缺失,导致“道德虚无”。人工智能在虚拟化社交平台中的应用往往出现责任逃避和道德选择等问题。例如,当虚拟角色在社交互动中引发纠纷时,承担责任的主体究竟是提供服务的社交平台、开发和应用人工智能的团队,还是承载算法的人工智能。这就容易导致人类将责任甩锅给人工智能,而人工智能究竟能不能承担起作为责任方的角色,学界一直以来也是争议不断。由于现有法律尚未明确人工智能的主体资格,无法有效规范人工智能在虚拟化社交中传播虚假信息、实施情感操控甚至引发心理伤害行为的行为。由此可见,责任主体的边界被模糊化了,伦理规范的滞后性一定程度上制约了社交平台的健康发展。再者,责任主体的模糊也就加剧了“道德虚无”的产生。由于责任主体的不明确,人工智能应用产生的不良后果如果无法得到明确追责,同时,不同国家和地区在虚拟社交伦理标准上存在差异,一些全球化运营的社交平台还面临着跨文化伦理冲突。久而久之,人类的道德底线将会被突破,道德遭受前所未有的挑战,甚至不复存在,人们逐渐变得麻木,从而导致“道德虚无”的产生。

4. 人工智能在虚拟化社交平台中应用困境的化解路径

时至今日,人工智能发展过程中所面临的困境几乎全部显现出来,过度张扬的技术理性产生了一系列的伦理问题,我们必须重新审视人工智能的发展和未来,以寻求人工智能与人类自身命运的和谐共存。

4.1. 实现技术突破和更新

虚拟化社交平台中,数据收集、存储、传输和使用的各个环节都可能存在数据和隐私泄露的风险,应更新保护和储存数据隐私的新的技术形式。可以考虑在人工智能对“隐私保护”伦理的内建过程中设置“保密与保密例外”原则,即人工智能对于筛分的信息数据中涉及个人隐私的信息采取保密、保护措施,而对于那些不涉及个人隐私的或违反法律规定的信息,采取保密例外处理[3]。并允许数据在加密状态下进行计算,通过动态身份认证和减少访问权限,降低数据泄露风险。虚拟化社交平台中充斥着大量无效信息和虚假内容,要开发能够筛选和过滤这些内容的技术。可以反过来利用人工智能的特点,开发能够有效识别和过滤虚假信息、不良内容的算法和工具,从源头遏制虚假信息和不良内容的传播。提高算法的透明度与可解释性,基于因果推理理论,将算法决策逻辑转化为可视化的因果关系图,使复杂的决策过程变得清晰可理解。同时,构建交互式的算法解释界面,用户可自主查询决策理由,监管机构也能对算法运行进行全过程监督,以增强公众对算法的信任。

4.2. 完善相关法律法规和管理制度

为了更好地规范人工智能在虚拟化社交平台中的应用，必须完善相关法律法规。通过制定清楚明晰的法律法规条文，界定人工智能开发团队、社交平台运营者和用户等不同主体的权利，明确各自的义务和责任，让用户实现有法可依。完善有关于人工智能开发与应用的政策法规，让开发者对算法透明性负责，保障用户对自己相关数据隐私的拥有权和知情权，并规范管理自身行为，不利用用户的数据隐私进行非法牟利。有关部门应加强执法力度，构建跨部门、多领域协同的执法监管机制，对于隐私侵权行为，加大惩处力度，形成法律威慑。还要加强执法监管，通过建立健全执法监管机制，加强对社交平台中人工智能应用的监督检查，严厉打击传播虚假信息、侵犯隐私和算法歧视等违法违规行为，维护良好的网络社交秩序。此外，虚拟化社交平台应加强平台管理，建立快速响应机制，及时处理有关问题，应建立健全内容管理制度，加强对平台内容的审核和管理，规范用户行为，积极防范和化解伦理风险。

4.3. 构建伦理准则体系

要从责任伦理的角度对人工智能的应用进行伦理规制并构建道德规范。责任伦理是约纳斯看到技术发展对人类的反噬而寻求人类自我保护的一套理论，责任伦理反对一切形式的狂热，不论是对技术的使用还是对自然的开发[5]。因此，为更好的抵制人工智能发展带来的伦理困境，要形成制度规定、价值导向等促进人工智能向善发展的内容。例如可以采取“蓄水池式”的公共决策机制，力求将人工智能发展与应用放置在人的意图之下[6]。还要加强伦理审查与评估。设立单独的人工智能伦理审查委员会或专业评估机构，吸纳伦理学、科学技术和法律等多学科专家及用户代表共同参与。人工智能在虚拟化社交平台中的应用，要对其功能设计、数据处理流程和潜在社会影响等进行全面伦理审查，当然也要持续跟踪评估，定期分析用户反馈与运行数据。若发现人工智能引发负面问题，要及时介入纠正，通过优化算法和调整功能设置等方式，维护虚拟化社交平台健康有序发展。

4.4. 开展相关教育活动

有关部门可以通过媒体宣传向公众科普人工智能在社交平台中的应用边界与治理措施，也可以邀请专家解读人工智能应用的相关原理和潜在风险，实现对公众的教育，引导社会公众正确认识人工智能应用的优势和风险，推动人工智能合理应用。对虚拟化社交平台中的不同用户进行分类教育。针对普通用户来说，通过社交平台的弹窗提示或者新手教程引导的方式普及人工智能在社交平台中的具体应用，明确告知其使用规范，让他们理清虚拟身份与现实生活的联系，提升辨别和自我保护的能力。对于青少年则主要是引导他们树立正确的网络社交价值观，可以联合学校开设人工智能素养课程，帮助他们在社交平台中正确使用人工智能。对于社交平台中的内容创作者，应加强教育提升其法律法规意识，遏制其利用人工智能虚构内容并大肆传播。通过持续性教育引导，实现用户从被动接受者向主动治理者的角色转变，共同构建安全健康的虚拟化社交生态。

5. 结语

人工智能的发展有着巨大潜能，当前各种 AI 软件的风靡就是最好的例证，我们已经走上了人工智能发展这条道路，未来将会出现更多的强人工智能产品，人工智能的发展只会日益深入。本文通过对人工智能在虚拟化社交平台中应用现状和形式的分析与梳理，总结概括了其中存在的困境，并探索出了能够有效解决这些问题的化解路径，为人工智能在虚拟化社交平台中应用的健康发展提供了理论范式，具有一定的参考借鉴意义和实践价值。当然，本文也存在一定的不足，例如，本文研究领域相对局限，对于人工智能在其他领域的应用没有进行更深入的探讨和分析，未来将对人工智能在其他领域中的应用和发

展出现的新伦理问题以及不同文化背景下人工智能的伦理差异等问题进行后续研究。我们必须意识到，在未来人工智能更进一步发展而威胁到人类自身的生命存在之前，人类应该尽早积极采取相关防范措施，解决问题并时刻保持警惕，只有这样，人类才不会在科学技术席卷的浪潮中迷失自我。

参考文献

- [1] 段伟文. 人工智能时代的价值审度与伦理调适[J]. 中国人民大学学报, 2017(6): 98-108.
- [2] 约翰·塞尔. 心、脑与科学[M]. 杨音莱, 译. 上海: 上海译文出版社, 2006.
- [3] 冯永刚, 席宇晴. 人工智能的伦理风险及其规制[J]. 河北学刊, 2023(3): 60-68.
- [4] 2022 主流社交媒体平台趋势洞察报告[EB/OL]. https://www.sohu.com/a/565274469_121101099, 2025-06-01.
- [5] 王玉雯. 人工智能技术的伦理审视[D]: [硕士学位论文]. 苏州: 苏州科技大学, 2021.
- [6] 田广兰, 单杰. 科技前沿的伦理挑战——第十一次全国应用伦理学研讨会综述[J]. 哲学动态, 2019(3): 127-128.