

反思与构建：生成式人工智能(AIGC)的伦理风险及其治理路径

马肖华, 王 宁

中原工学院马克思主义学院, 河南 郑州

收稿日期: 2026年7月3日; 录用日期: 2026年7月25日; 发布日期: 2026年8月5日

摘 要

以ChatGPT为代表的生成式人工智能(AIGC)掀起了新一波人工智能热潮。AIGC因其可以在算法的帮助下产出如文字、图片、音视频等内容, 为各行各业带来无限创新的机会而引起人们的巨大关注, 并迅速成为热议话题。作为最具代表性的颠覆性技术, AIGC在技术红利释放的同时, 也给人类生活带来了如人的主体性消解、信息安全、隐私侵犯、算法歧视等新的风险挑战, 引发重大的伦理关切。应对AIGC的伦理风险, 可通过制定正确的科技发展方向、构建技术风险防范监管治理共同体、出台相关伦理法律规范, 完善其伦理治理体系, 共同致力于科技向善。

关键词

AIGC, 技术伦理, 伦理治理

Reflection and Construction: The Ethical Risks of Generative Artificial Intelligence (AIGC) and Its Governance Path

Xiaohua Ma, Ning Wang

School of Marxism, Zhongyuan University of Technology, Zhengzhou Henan

Received: July 3, 2026; accepted: July 25, 2026; published: August 5, 2026

Abstract

The generative Artificial Intelligence (AIGC) represented by ChatGPT has set off a new wave of artificial intelligence craze. AIGC has attracted great attention because it can produce content such as

text, pictures, audio and video with the help of algorithms, bringing unlimited innovation opportunities to all walks of life and quickly becoming a hot topic. As the most representative disruptive technology, AIGC has brought new risks and challenges to human life, such as the elimination of human subjectivity, information security, privacy invasion, algorithm discrimination, and so on, while releasing technological dividends, causing major ethical concerns. To deal with the ethical risks of AIGC, we can improve its ethical governance system by formulating the correct direction of science and technology development, building a community of technical risk prevention, supervision and governance, and issuing relevant ethical laws and regulations, so as to jointly strive for the good of science and technology.

Keywords

AIGC, Technical Ethics, Ethical Governance

Copyright © 2026 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

AIGC (全称 AI generated content), 即生成式人工智能, 指在人工智能算法帮助下创建内容的技术, 其强大的内容创作能力迅速引起人们的巨大关注。2022 年末, OpenAI 推出了人工智能聊天机器人 ChatGPT (Chat Generative Pre-trained Transformer), 作为 AIGC 的一个典型产品, ChatGPT 聚焦于文本生成。相较于此前的人工智能产品, 该模型在代码理解及内容生成能力上取得了显著进展, 在文字交互领域表现卓越, 使人工智能这一领域再次成为人们的关注热点。随着人工智能算法学习能力的提升, AIGC 技术快速发展, 赋能各类场景, 在全球范围内展示了巨大的潜力。然而, 任何技术都是一把“双刃剑”, 我们也必须认识到技术发展过程中势必会带来一定的风险挑战, 在继续拓展 AIGC 技术的应用领域时, 需正视其所带来的社会风险、伦理风险, 在创新与安全、发展与伦理之间寻求平衡, 以科技向善引领 AIGC 技术创新发展和应用, 让人工智能技术更好造福人类社会。

2. AIGC: 内容生成新范式

内容生产与传播模式主要有专业生产内容(PGC)、用户生产内容(UGC)、专业用户生成内容(PUGC)等, 这几种模式都是以人为主体的生产内容, 但 AIGC 的出现却是开启了以机器为主体的内容生产模式, 内容生产产生了大变革[1]。梳理 AIGC 的发展历程, 大概可以划分为三个阶段。

第一阶段, 20 世纪 50 年代到 20 世纪 90 年代初期, 此阶段是 AIGC 早期萌芽阶段, 人们开始探索利用人工智能技术生成内容, 但此时技术水平相对有限, AIGC 仅限于小范围的研究试验, 并且大多是通过预先设定的算法来生成内容[2]。第二阶段, 20 世纪 90 年代初期到 2010 年, 伴随着深度学习算法的提出, AIGC 的智能及感知能力得到了进一步提升[3]。第三阶段, 2010 年至今是 AIGC 的迅速发展阶段, 此为其黄金发展阶段, 生成式对抗网络、生成式预训练模型等算法的提出和迭代更新, 为 AIGC 提供了强大的技术支撑, 使 AIGC 在自然语言处理、图片和音频生成等领域取得了重大突破, 智能能力由算力向感知拓展。同时, 该技术在产业端的应用也呈现出百花齐放样态, 微软、英伟达、DeepMind、OpenAI 等科技公司纷纷推出了各类应用, 如 2021 年 OpenAI 推出了 DELL-E 模型, 一年后又推出了升级版 DELL-E2, 主要应用于文本和图像的交互, 近期又推出了最新版本 DELL-E3 [4]。

与传统 AI 相比, AIGC 的显著特征主要表现在以下三个方面。首先是自主性。AIGC 的核心是创造, 其通过从数据中提取、学习要素, 在算法与算力的加持下, 无需人工干预或手动调整就能够根据用户需求或特定任务要求自动地生成和调整内容。由于 AIGC 是基于深度学习技术进行训练和推断, 因此它还能够根据用户的需求或任务要求, 自动地调整生成内容的质量和风格以满足不同的需求和要求。这种自主性使 AIGC 在许多应用场景下比其他人工智能技术更加高效和灵活, 同时也在很大程度上提高了生产力水平。其次是普遍性。AIGC 已广泛应用于教育、影视、医疗等领域, 各行业多元化场景应用落地。不同行业可以根据自身的需求和特点, 利用 AIGC 解决特定的问题或提高效率, 例如虚拟人直播、智能助手等。它具有比较简单的操作界面, 这使得用户可以轻松上手并利用该技术进行创作、分析或解决问题, 这种对用户的友好性让更多人可以接触到并使用该技术, AIGC 在普通用户中得到了广泛的普及和应用。最后是简易性。AIGC 的操作通常基于自然语言处理技术, 用户不需要通过漫长的学习, 也不需要具备相关学科背景和专业知识, 只需通过输入自然语言指令或文本, 告诉 AIGC 需要生成的内容即可获得想要的答案。

总的来说, AIGC 经历了从早期探索到统计学习方法的发展, 再到深度学习和大规模预训练模型的兴起, 随着技术的不断进步, 其已经被广泛应用于游戏设计、广告创意、新闻报道、新媒体运营等多个应用领域, 并展现出了巨大的潜力, 并且在未来有望继续创新和发展。

3. AIGC 的应用失范与伦理风险

美国技术哲学家詹姆斯·摩尔曾经提出: 伴随着技术革命, 社会影响增大, 伦理问题也会随之增加。AIGC 的发展和应用会带来社会生产关系、生活方式、组织结构和运行状态的联动变革, 在变革中, 新的社会风险也将随之产生[5]。

3.1. 人的主体性消解: 智能技术异化人机关系

正因人们对媒介(技术)如何影响潜意识抱温顺接受的态度, 才使得媒介(技术)成为囚禁其使用者的无墙的监狱[6]。随着人工智能表现出来的越来越智能化, 社会上也出现了很多关于 AIGC 是否具有自我意识问题的讨论, 由此担心人的主体地位是否会发生变化、人们是否会失去创造力等问题。语言和思维是人类区别于其他动物的独特标志, ChatGPT 的出现使主客间变得模糊, 其所呈现的主体性特征, 正在弱化人的主体性地位。以往的技术改造自然, 虽有不良副作用, 并非不可拯救, 今天的技术却准备改造或重新定义人的存在, 这意味着人类主体正在变成客体。这是自语言的发明使人成为主体以来的主体客体化大变局[7]。首先, ChatGPT 逐渐掌握人类语言处理能力, 在模拟人类对话时, 将自身置于情境之中并生成与人类共鸣的情感内容, 展现出类似人类的情感反应, 这让我们在语言、思考等方面对技术的依赖增长, 这可能会造成自我个性化的丧失, 让人的属性逐渐被技术所淹没。其次, AIGC 技术可以自动处理和分析大量的文本数据, 从而对语言进行规范化和标准化, 这种规范化可能导致语言的多样性和灵活性降低, 限制了人类使用语言的自由度和创造力。其通过算法和模型来模拟人类的思维过程, 可能导致人类的思维模式逐渐趋于一致, 失去独立思考和创新的能力, 这使得人们容易产生惰性。当人们过度依赖人工智能的解决方案时, 不再主动思考问题, 可能会失去解决问题的能力, 从而在面对新问题时变得无所适从。同时, 人们习惯于使用 AIGC 后, 可能会更加倾向于使用简化和机械化的语言表达方式, 这种表达方式会使人们失去对自己语言能力的掌控和精细化的表达能力, 导致语言的贫乏和单调。

3.2. 信息安全挑战: 信息洪流冲击用户安全

安全问题始终是 AI 技术发展和应用中不可回避的难题。随着 AIGC 技术的发展, 安全问题变得愈加严峻, 主要表现在三个方面。首先是产生了大量的虚假信息。AIGC 凭借其强大的自主性和灵活性, 对信

息生成与传播的贡献达到了前所未有的高度,但它是基于对已有数据的学习,在一些问题上可能会出现“一本正经地胡说八道”。从技术的本质上来看,它所生产的内容是在海量信息中筛选出来的内容,因此 AIGC 自身无法判断所生产的内容是否是真实准确的,这就导致 AIGC 产生大量虚假信息,部分用户利用 AIGC 生产政治敏感性等有害信息,这会对人们的思维和行为产生误导。尤其是 ChatGPT 的出现,文字的高信息量及个人化的、实时性的聊天方式都极大降低了用户对虚假有害信息的辨别和屏蔽能力,扩大了相关信息的传播和覆盖面,进而以低廉的成本和易滋生的特性吸引更多蠢蠢欲动的谣言散播者的加入[8]。其次是诱发新型违法犯罪。不法分子利用 AIGC 图片生成、语音合成等技术假冒他人身份进行诈骗、诽谤等违法行为。当前,已经出现利用 AIGC 实施犯罪的案例。随着 AIGC 的日益成熟,该技术甚至可能会被利用制造虚假政治信息,操控人们的舆论从而给国家安全带来巨大威胁。最后是数据泄露。AIGC 在内容生成过程中需要海量已有数据训练,公司可以通过多途径获取用户信息,这极有可能触及到客户的隐私数据。以 ChatGPT 为例,在 ChatGPT-4 最新发布之际,有用户表示在社交媒体上看到了其他人的历史搜索记录标题,其拥有超过 1750 亿的参数,一旦发生数据泄露事件,将造成难以想象的后果。

3.3. 算法歧视浮现: 技术系统承载价值倾向

在人工智能大背景下,考虑数据安全问题之余,算法的不可控性所带来的伦理问题也是需要引起人们关注的核心。首先, AIGC 的算法主体嵌入过程,是一种技术植入及主导内容生成的过程,工程师能够通过操控算法的方式将自己的偏好植入训练数据中,进而引导文本输出特定价值观,如果工程师持有历史错误、文化偏见或种族歧视的价值观,这些观念或将在模型与用户交互中,潜在地传播不良意识形态,从而形成算法歧视。如 AIGC 在图像生成方面就存在较为显著的歧视现象。其次,技术的本质是能力,而能力越大,其博弈均衡点就对技术掌权者越有利。如果技术掌握在少数人手里,弱者的讨价还价收益就越小。在全球各种思潮、文化和价值观念的相互碰撞下, AIGC 技术很有可能被政治操控用作意识形态宣传,掌握大数据核心技术的西方发达国家,按照其价值观制定全球政治秩序和规则,裁剪符合其意识形态标准的数据库,使得全球信息体系和政治秩序中既存的不平等和垄断现象被进一步强化,形成数据强权[9]。

3.4. 责任判定风险: “机器权力” 贬斥 “个体自由”

AIGC 的快速发展及广泛应用也带来了相关的责任判定方面的争议。首先是主体责任模糊。ChatGPT 等生成式人工智能的应用模糊了负责任创新与研究的主体责任界限[10]。同时, AIGC 在使用过程中难免会帮助人们做出各种决策或者选择,当这些决策或者选择违背社会秩序和法律时,谁要为这些决策或者选择负责任,法律上还没有明确规定。AIGC 与人不同,由于其缺乏肉身性,因而在身体体验(即感受性层面)就缺乏一种海德格尔语境中对死亡的“畏”,如果没有对“惩罚”的恐惧,“责任”意识也无从谈起,“身体”的存在是“责任”机制运行的根本前提[11]。因此 AIGC 本身不会成为责任主体,但是工程师赋予了 AIGC 创作能力,它所辅助人们所做的决策难免不会体现工程师的价值观,而且,使用者用该技术帮助自己做决策,他们也是技术的操纵者,这种主体责任的模糊就导致了权责归属不明确,追责更加困难。其次是道德边界模糊。AIGC 缺乏人类所固有的社会意识和道德判断力,容易导致 AIGC 在处理涉及道德和伦理的复杂情境时无法像人类深入考量伦理道德原则以及人文关怀[12]。以 ChatGPT 为例,当它在回答人们所问问题时会受到人们态度的影响,做出难以前后一致的道德决策。ChatGPT 是以海量数据为基础进行训练,在此过程中,语料库中的部分有毒信息会进入到生成内容中,当 ChatGPT 在模糊的道德边界上游离时,会误导使用者的选择倾向,在公共决策过程中,可能会忽视部分群体,从而导致 ChatGPT 做出功利主义决策。

4. AIGC 伦理风险成因分析

AIGC 技术在全球范围内的蓬勃发展展示了巨大的潜力, 为各行各业带来无数创新机会。与此同时, 我们也必须认识到在其发展潜力无限的同时也面临更多不可忽视的挑战。

4.1. 迭代的风险: 技术存在局限性

技术的发展并不是一帆风顺的, 会呈现出螺旋式上升、波浪式前进的趋势, 因此其中必然会有不足之处。从算法方面看, “算法黑箱”导致的不透明性与不可解释性为技术人员带来了挑战, 现有技术尚未完善, 难免出现算法偏见等问题。以 ChatGPT 为例, 首先是文本数据自身存在偏见, RLHF 是一种结合了强化学习和人类反馈的智能模型, 这种模型将人类视为智能系统的一部分, 通过与人类合作来优化决策过程, 但其本身嵌入的算法不能鉴别文本数据的价值偏见, 会导致带有偏见的内容输出。在 RLHF 中, 算法的决策依赖于训练数据, 这些训练数据来自于不同的人所产出, 每个人对同一事件的考虑不同, 所做出的结果也不同, 因此训练数据本身也就存在着价值偏见, 算法可能会学习到这些偏见, 那么算法决策就会相应的呈现偏见, 例如, 如果训练数据中白人的表现比其他人更好, 那么算法可能会倾向于预测白人更优秀。其次, 输出文本也受到工程师的主观偏见, 因此 ChatGPT 在进行伦理价值判断时往往会做出多种选择。最后, 使用者也会有自己的偏见, 这就导致多次询问相同或相似问题时, 结果常呈现显著差异, 有时甚至会提供模棱两可的答案。

4.2. 把关的重申: 主体伦理素养淡薄

AIGC 主体责任意识淡薄可以说是引起技术伦理风险的又一大隐患。首先是工程师道德责任感的缺失。汉斯·尤纳斯的“责任伦理”思想, 来源于对技术时代人类发展的忧思, 他提出的责任伦理即“如此行动, 以便你的行为后果与人类持久的真正生活一致”, 也就是说人类有责任确保行为能够实现人的持久生存[13]。人工智能在研发的过程中工程师占有绝对的主导权, 研发过程中也势必会受到外界环境的影响, 一些工程师为确保研究顺利推进, 有时受经济、政治等外界因素所迫, 会做出违背伦理道德乃至社会规范的行为。或者工程师过于关注技术的先进性和创新性, 不顾技术对社会的不利因素和潜在风险, 这会导致技术向“恶”的方向发展, 不利于人类社会长期可持续发展。其次是 AIGC 使用者伦理素养不高。一方面, 一些使用者缺乏对 AIGC 技术的深入了解和认识, 不了解其潜在的风险和责任, 他们可能没有意识到自己的行为可能对他人或社会产生负面影响, 因此缺乏对此类行为的重视和责任感。在使用过程中可能出现失误或不当行为, 这些行为可能导致数据泄露、系统崩溃或其他不良后果, 但使用者可能没有意识到这些问题的严重性。另一方面, 一些使用者可能只是利用 AIGC 完成任务或者是获得结果, 这就会导致他们对技术过于依赖, 丧失自身的主动性, 失去创造力, 甚至可能为了追求个人或组织的利益, 忽视了道德和伦理规范, 将利益置于首位。

4.3. 制度的缺位: 伦理法律规范滞后

科技的发展速度非常快, 新技术层出不穷, 技术的发展始终存在这样一个悖论: 我们永远不可能为一项新出现的技术做好充分准备, 包括法律、规则和反制技术; 但如果我们一味踌躇于有没有做好准备, 反而又会限制技术的发展, 陷入竞争劣势。对于已经到来的 AIGC 大潮, 我们还面临着众多的不确定性, 但随着 AIGC 社会影响增大, 其发展超越了已有的伦理边界, 伦理法律规范的滞后使得带来的伦理风险也在不断增加。首先是科技伦理审查制度的不完善。目前尚未形成完善的伦理原则和监管制度去把控技术的研发、使用过程。其次是相关法律法规规范滞后。就目前来看, 人们大多追求短期利益, 在制定相关法律法规时或许并未将人类长远的效益考虑进去, 这就导致制定的相关政策与技术发展不匹配。如, AIGC

融入人们的社会生活后已经引起了一些权责纠纷问题, 但缺乏相应的法律依据、可操作性的法律条例, 当前的法律制度无法应对相应风险挑战。虽然我国于 2023 年出台了《生成式人工智能服务管理暂行办法》¹, 但是该法规仍存在一些不足之处, 比如缺乏专门预防算法偏见的具体规定、未对 AIGC 使用范围做出明确的规定等。

5. 应对 AIGC 伦理风险的治理对策

AIGC 技术在全球范围内的蓬勃发展展示了巨大的潜力, 为各行各业带来无数创新机会。与此同时, 我们也必须认识到在其发展潜力无限的同时也面临更多不可忽视的挑战。

5.1. 技术的完善: 制定正确的科技发展方向

维贝克的“道德物化”思想提示我们, 技术并不是任由人使用的、价值中立的工具, 而是与人互相建构的, 共同完成特定的能动行为。AIGC 虽展现了强大的内容创作能力, 但其技术的发展尚不够成熟, 后期发展难免遇到瓶颈, 应该进一步完善技术自身, 努力将特定的公共价值嵌入技术, 确保朝向正确的发展方向。首先, AIGC 在发展过程中应该始终秉持以人为本的理论原则。马克思认为, 科技是为人类服务的, 自然科学通过工业日益在实践上进入人的生活, 改造人的生活[14]。因此脱离 AIGC 对人的束缚, 正确认识技术是辅助人类的工具, 维护人类自身尊严, 坚持“以人为本”, 以实现人的自由全面发展作为最终目标。其次, 应加强开发行业技术标准和规范, 同时探索技术手段, 助力技术监管, 制定和完善生成式人工智能技术的技术标准和规范, 确保技术和产品的质量 and 可靠性。从技术自身出发, 利用更先进的算法技术提高模型的性能和泛化能力。在预训练阶段, 增强训练数据的质量和多样性, 清晰数据收集标准, 广泛收集不同国家的数据, 合法合规将同意使用的数据纳入模型中, 保持数据库不断更新, 提升 AIGC 理解其他国家语言的准确性, 避免 AIGC 技术实行数据“垄断”。此外, 可以利用可视化技术增加 AIGC 的可解释性和透明度, 帮助使用者更好地理解模型的决策依据和输出结果。用户对 AIGC 的可解释性需求主要包括透明性、解释性、可信度和可追溯, 满足这些需求可以增加用户对人工智能系统的信任和接受度, 改善用户体验[15]。同时, 正确看待我国与西方技术的差距, 自立自强, 加快研发独立于西方技术的生成式人工智能系统, 建立符合正确价值导向的数据基础。虽然 ChatGPT 是保持“价值中立”原则, 但是仍然有瑕疵, 只有拥有国产自主 ChatGPT, 才能规避技术被西方阻碍等问题, 用适合国家发展的数据体系规范算法, 加速形成以正确价值引导的话语体系, 彰显技术价值, 做好思想工作。最后, 面对 AIGC 的潜在技术风险, 可以沿着技术反制、技术对抗等技术创新的思路应对, 以降低技术带来的风险挑战。比如提高 AI 技术的透明度, 通过增强 AI 系统的可解释性, 可以更好地理解 AI 的决策过程和结果, 从而更好地判断其是否被滥用。开发检测和识别算法, 识别和过滤机器生成的内容, 或者利用自然语言处理技术, 对机器生成的内容进行纠正和改进。

5.2. 体系的建立: 构建技术风险防范管理共同体

打造生成式人工智能伦理风险防范管理共同体, 促进多方协同共治。生成式人工智能的伦理风险涉及面广, 复杂程度高, 风险治理涉及智能技术、伦理哲学、公共管理等多领域, 需要政府、企业、高校、行业组织、公众等多主体搭建多方协同共治的组织网络, 明确不同主体的权责义务, 建构多主体合作治理的紧密关系, 形成集群治理效应。为了应对 AIGC 时代人机交互中存在的伦理风险, 政府应当进一步明确监管框架和法律法规, 建立合理的治理框架, 完善相关法律法规, 明确各方的责任和义务, 规范 AIGC 技术的使用和管理, 加强监管和评估, 确保技术的安全性和公正性。社会应建立恰当的人机交互伦理规

¹https://www.gov.cn/zhengce/202311/content_6917778.htm

范, 积极参与人工智能技术的监督和评估, 通过建立社会监督机制, 鼓励公众提出建议和反馈, 促进技术的公开透明和公正性。社会还可组织开展 AIGC 技术风险防范管理的宣传教育活动, 加强公众对该技术的认知, 提高伦理意识和风险意识。企业、工程师、平台等应采取措施将机器“道德化”, 帮助机器发展出合乎道德的语言表达能力。企业在 AIGC 技术研发、管理中扮演着重要的角色, 应强化风险管理, 严格遵循法律法规和伦理规范, 加强技术研发风险评估和伦理审查, 建立健全的内部管理机制, 避免技术的不当使用和滥用, 及时识别和处理潜在的安全漏洞和风险事件, 同时建立事件应对预案, 及时有效处理各类事件。工程师作为 AIGC 技术的创造者和开发者, 其职业道德与伦理规范对于保持行业的良好形象以及保证技术的稳定、可持续发展至关重要。工程师应严格遵守相关的伦理规范和法律法规, 在具备技术素养的同时, 要提高自身的伦理素养和判断能力, 更好地把握技术的伦理边界和责任, 在经济利益和社会责任之间做出平衡, 帮助机器发展出合乎道德的语言表达能力。使用者应当了解 AIGC 技术伦理边界和社会责任, 增强风险意识和安全意识, 避免技术的不当使用和滥用, 确保技术的安全性和有效性。还应当积极参与 AIGC 技术的监督和反馈机制, 提供对技术的评价和建议, 促进技术的改进和发展, 为构建技术风险防范管理共同体做出贡献。

5.3. 法规的兜底: 完善相应伦理法律法规

法律的强制性虽然可以将 AIGC 限制在一定范围之内, 但随着当下技术的进步与日俱新, 立法可能在及时性上有所欠缺, 因此需要“道德”作为有效补充手段, 综合“道德”以及“法律”二者的约束, 使得 AIGC 张弛有度, 良性发展。完善 AIGC 的伦理准则和行业标准, 明确 AIGC 系统的设计、开发和应用应遵循的道德原则, 同时对 AIGC 系统的性能和安全性进行严格规范。行业组织可以制定可信的 AIGC 伦理指南, 以更好支持 AIGC 技术的健康可持续发展。建立 AIGC 伦理准则和行业标准, 明确 AIGC 技术在开发、应用时应遵循的道德原则, 能确保 AIGC 系统的可靠性和可持续性, 为促进 AIGC 的健康发展提供有力保障。面对 AIGC 技术的不确定性, 建立伦理风险评估机制, 梳理风险源头、风险类型、风险强度, 思考风险产生原因、风险解决措施, 防范未发生的风险, 对 AIGC 应用的负面影响开展预见性风险评估, 提前做好风险预案。

近年来, 互联网迅速且极大地改变了人们的生活样态, 相对滞后的法律和规则一直在后面气喘吁吁地追赶。截至目前我国已经制定出台网络领域立法 140 余部, 基本形成了较为健全的网络法律体系。但是由于新情况新业态层出不穷, 特别是 AIGC 等人工智能的突飞猛进, 又给网络法治带来了新一轮的冲击。AIGC 技术具有超强的进化速度和迭代能力, 与之相关的法律法规却并未能够实现同步跟进, 因此法律和监管也应求变, 可以在法律制度上进行有针对性的、超前性的设计, 实现科技创新与规则治理的双轮驱动, 避免滞后性。各国对此纷纷出台有关举措。2023 年 7 月, 《生成式人工智能服务管理暂行办法》²正式对外公布, 成为我国生成式人工智能产业首份监管文件, 管理办法将重点解决当前阶段境内外生成式人工智能技术“狂飙”途中的热点问题。法律作为保障 AIGC 有所为、有所不为的重要后盾, 也是引导 AIGC 良性发展的必要举措, 因此, 应该结合我国科技发展现状, 制定一套合理的法律法规来约束技术的规范发展。比如, AIGC 出现生成内容的虚假陈述、知识产权侵权与版权问题, 需要就试点阶段的知识合法性与共享问题设置法律底线, 明确人机合作生成内容的权益与权责分配法规[16]。还可以通过设定数据保护、隐私保护等法律条款, 确保 AIGC 技术的健康、稳定发展, 真正造福人类社会。

6. 结语

AIGC 技术在全球范围内的蓬勃发展展示了巨大的潜力, 为各行各业带来无数创新机会。与此同时,

²同脚注 1

我们也必须认识到在其发展潜力无限的同时也面临更多不可忽视的挑战。

基金项目

河南省软科学研究项目“中国式现代化视域下生成式人工智能(AIGC)的伦理问题及其治理研究”(242400411117)阶段性成果。

参考文献

- [1] 郭全中, 袁柏林. AIGC 与 WEB3.0 有机融合: 元宇宙内容生产的新范式[J]. 南方传媒研究, 2023(1): 36-47.
- [2] 葛智君, 刘梦源, 罗剑武. AIGC 的发展与挑战综述[J]. 电子产品可靠性与环境试验, 2025, 43(1): 114-123.
- [3] 王泽轩, 陈亚军. AIGC 技术发展与应用进展[J]. 印刷与数字媒体技术研究, 2024(4): 1-14+96.
- [4] 中国信息通信研究院, 京东探索研究院. 人工智能生成内容(AIGC)白皮书(2022 年) [R]. <http://www.caict.ac.cn/kxyj/qwfb/bps/202209/P020220902534520798735.pdf>, 2024-09-02.
- [5] 王文玉. 生成式人工智能应用的社会风险与治理路径[J]. 南京邮电大学学报(社会科学版), 2024, 26(5): 28-39.
- [6] 埃里克·麦克卢汉, 弗兰克·秦格龙. 麦克卢汉精粹[M]. 第2版. 何道宽, 译. 北京: 中国大百科全书出版社, 2021: 211.
- [7] 赵汀阳. 人工智能提出了什么哲学问题? [J]. 文化纵横, 2020(1): 43-57.
- [8] 朱嘉珺. 生成式人工智能虚假有害信息规制的挑战与应对——以 ChatGPT 的应用为引[J]. 比较法研究, 2023(5): 34-54.
- [9] 郑元景. 大数据环境下我国意识形态安全风险与治理策略[J]. 中国社会科学院研究生院学报, 2016(5): 17-22.
- [10] 刘瑶瑶, 梁永霞, 李正风. 生成式人工智能与科研伦理: 变革、挑战与展望[J]. 科学观察, 2024, 19(4): 1-8.
- [11] 张卫, 黄成驰. 生成式人工智能中的“责任”问题及成因探析——基于身体的视角[J]. 图书馆建设, 2023(4): 15-18.
- [12] 解学芳, 曲晨. 价值对齐: AIGC 时代的人工智能文化科技伦理风险与精准共治路径研究[J]. 兰州大学学报(社会科学版), 2024, 52(3): 147-156.
- [13] 周英豪, 石诚. 人类增强技术的伦理问题及其治理——基于汉斯·约纳斯责任伦理的视角[J]. 佛山科学技术学院学报(社会科学版), 2024, 42(3): 38-44+58.
- [14] [德]卡尔·马克思. 1844 年经济学哲学手稿[M]. 北京: 人民出版社, 2000: 89.
- [15] 陈建兵, 王明. 负责任的人工智能: 技术伦理危机下 AIGC 的治理基点[J]. 西安交通大学学报(社会科学版), 2024, 44(1): 111-120.
- [16] 俞鼎, 李正风. 生成式人工智能社会实验的伦理问题及治理[J]. 科学学研究, 2024, 42(1): 3-9.