

生成式人工智能侵权的法律规制研究

常启龙

西南林业大学文法学院, 云南 昆明

收稿日期: 2026年7月12日; 录用日期: 2026年8月11日; 发布日期: 2026年8月19日

摘要

生成式人工智能的深度应用在推动数字经济发展的同时, 也对传统侵权规制体系带来了新的挑战。本文系统梳理了生成式人工智能在数据训练阶段的著作权侵权、内容生成过程中的人格权侵害以及深度伪造技术衍生的多元危害, 并对当前法律规制在责任主体认定、因果关系证明及归责原则适用等方面面临的多重困境进行分析。在结合我国产业发展需求的基础上, 提出构建多层次归责体系、完善因果关系认证的技术化路径、引入公证证据保全机制及推动多元主体协同治理等建议, 以期为技术创新与权益保护的平衡提供理论支撑与参考。

关键词

生成式人工智能, 侵权风险, 法律规制

Research on Legal Regulation of Generative Artificial Intelligence Infringement

Qilong Chang

School of Humanities and Law, Southwest Forestry University, Kunming Yunnan

Received: July 12, 2026; accepted: August 11, 2026; published: August 19, 2026

Abstract

Generative AI, while advancing the digital economy, poses new challenges to traditional tort liability systems. This paper reviews copyright infringements in data training, personality rights violations in content generation, and harms from deepfakes, and analyzes dilemmas in identifying liable parties, proving causation, and applying imputation rules. Considering China's industrial needs, it proposes multi-tiered imputation, technical causation verification, notarized evidence preservation, and multi-stakeholder governance, aiming to balance innovation and rights protection.

Keywords

Generative Artificial Intelligence, Infringement Risk, Legal Regulation

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. 生成式人工智能的侵权类型及特征

1.1. 生成式人工智能概述

生成式人工智能(Generative Artificial Intelligence, 简称 GAI)是人工智能技术在发展和应用过程中形成的一个重要分支,是基于算法、模型、规则生成文本、图片、声音、视频、代码等技术。在实际应用中,生成式人工智能能够依托多模态基础大模型,利用用户输入的相关资料,生成具有一定逻辑性和连贯性的内容,比传统的人工智能更进一步的是,生成式人工智能不仅能够对输入数据进行处理,更能学习和模拟事物内在规律自主创造出新的内容。当下生成式人工智能已经广泛应用到了工业生产、文娱创作、教育教学等多个领域并展现出了丰富的应用场景与潜力,如在工业生产中,生成式人工智能能够依托算法显著提升生产效率,在预测设备故障以及辅助设计迭代等方面也展现出巨大潜力;在文娱创作中,生成式人工智能可以应用于文本、图像、视频、代码以及音频生成等,对文化产品的生产提供技术支持;在教育教学中,生成式人工智能可以通过虚拟现实等技术进行模拟实验和情景教学,为教育教学打开了新思路。

然而,生成式人工智能的发展和应用也同样给传统法律规则体系带来了挑战。从法律属性的定位来分析,生成式人工智能并没有独立的行为能力和权利能力,本质上是一种为人所研发和使用的法律客体。2023 年颁布的《生成式人工智能服务管理暂行办法》将其定性为一种“服务”¹,其所有的行为均以人提供的基础材料与下达的操作指令为基础,这一属性定位是明确其侵权行为与责任和梳理构建规制体系的前提,也是生成式人工智能与传统民事主体相区分的核心依据[1]。

1.2. 生成式人工智能的侵权类型

生成式人工智能引发的侵权问题以前端数据训练中的著作权侵权、作品生成过程中的人格权侵权以及深度伪造技术衍生出的多元侵权三个领域为重点。

生成式人工智能并不能凭空产出作品,其要以前置的算法开发训练为基础,结合操作者提供的素材资料 and 具体指令来生成作品。前置的算法开发训练过程又包含数据输入、模型学习、模型优化阶段三个阶段,在数据输入阶段中,开发者往往会利用算法抓取受著作权保护的作品,进行转码、清洗等预处理,转码过程会涉及对作品的数字化复制,而数据清洗过程中又会删除原作者的姓名、版权声明等信息,在这个过程中往往会侵犯到原作者的复制权、署名权、保护作品完整权等著作权内容。

随着生成式人工智能在文娱创作中的不断应用,各类视频平台中的人脸替换、人脸生成、声音克隆等类别的视频层出不穷。依据《中华人民共和国民法典》(后文简称《民法典》)规定,肖像以社会一般人凭视觉识别特定自然人为标准²,也就是说在利用生成式人工智能进行人脸替换时,若去除面部特征后视

¹https://www.gov.cn/zhengce/zhengceku/202307/content_6891752.htm, 2026 年 7 月 17 日访问。

²参见《中华人民共和国民法典》第 1018 条及 1019 条。

频仍能凭发型或面部特征等使社会一般人识别出权利人,即侵犯了他人的肖像权,而当 AI 艺人当虚拟人脸撞脸现实人物时,若社会一般人能凭肉眼识别特定自然人也可构成肖像权侵权[2]。在未经允许的情况下利用生成式人工智能克隆配音演员或知名人物的声音用于视频配音的现象更是层出不穷,而依据《民法典》对声音进行保护的规定,这种行为无疑构成对他人人格权的侵犯³。

随着生成式人工智能技术的不断迭代,深度伪造技术被广泛应用于生成各类图像、音频与视频内容。这种技术不仅与传统的侵权模式完全不同,由于深度伪造技术的主要功能即是伪造,其技术设计天生具有侵权倾向。除侵害肖像权和声音权益外,深度伪造还极易引发名誉贬损、虚假信息指数级传播等复合型社会风险。例如,将某人面部转接至特定视频场景中,不仅侵害了当事人的肖像权,还进一步损害了当事人的名誉权,而且在责任主体认定上,深度伪造侵权往往涉及算法开发者、平台运营商及终端用户等多方主体,难以统一划分责任。

以侵犯肖像权为例,早在 2023 年上海市金山区人民法院就已经审理过类似案件,被告上海某公司在未经原告授权许可的情况下,私自将原告的视频中的形象利用生成式人工智能技术“换脸”后上架供其注册用户换脸使用,上海市金山区人民法院经审查认为,被告收集了包含原告人脸信息的出境视频,将该视频中的原告面部替换为自己提供的照片的面部,在合成过程将新的静态图片中的特征与原视频部分面部特征、表情等通过算法进行融合,生成新的视频内容,虽然视频内容与原告面部不尽相同,但仍能结合发型及其他特征识进行识别,系利用信息技术手段伪造等方式侵害原告的肖像权,应当认定侵害了原告的肖像权⁴。显然这与社会一般人能凭肉眼识别特定自然人也可构成肖像权侵权的法理基础是相契合的。

1.3. 生成式人工智能侵权的特征

与传统意义上的民事侵权行为相比较,生成式人工智能侵权呈现出了更具有时代性的特征,这也对传统的法律规制模式产生了新的挑战。

第一,侵权主体更加多元化且呈现链条式的分布。无论是数据训练端的著作权侵权,还是人脸替换中的人格权侵害,都涉及到了包括算法开发者、数据提供者、平台运营者与终端用户等多方主体[3]。以深度伪造换脸视频为例,算法训练设计、用户下达的指令与平台未删内容互相结合形成了链条式的侵权,传统意义上单一侵权行为人的模式已经不再适用。

第二,侵权过程的高度技术性。生成式人工智能以深度学习模型为基础,其进行逻辑思考和输出内容的过程难以解释和理解,在声音克隆、肖像替换及训练数据抓取中,算法的自主性又导致侵权行为隐蔽,形成典型的“黑箱效应”,呈现出了高度的技术性[4]。

第三,损害结果传播和扩大的快速性。不同于传统的 PS 等人工操作,深度伪造与 AI 生成内容只需要用户下达指令即可自动完成,且生成内容借互联网可在短时间内快速传播,使得损害后果瞬间放大。

第四,法益侵害的复合性以及举证的困难性。一个生成内容常会同时侵犯著作权、肖像权、名誉权等多项内容,如利用生成式人工智能将某人面部转移至特定视频场景中,既侵犯了他人肖像又贬损了他人名誉。且对被侵权人而言,侵权行为往往难以被及时发现,在被发现时通常也已造成了严重的不良后果,又由于生成式人工智能“算法黑箱”的技术特征,被侵权人也难以有效掌握有效证据,在诉讼过程中往往处于不利地位。

³参见《中华人民共和国民法典》第 1023 条。

⁴参见上海市金山区人民法院(2023)沪 0116 民初 4101 号民事判决书。

2. 生成式人工智能侵权的法律规制困境

2.1. 责任主体难以认定

追究侵权责任以责任主体明确为前提，然而在实践中确定责任主体面临诸多困难。首先，生成式人工智能的法律定位并不清晰，目前我国法律体系中大致将责任主体限定为人类或拟制人格主体，而人工智能本身是否具备法律人格尚无定论，在弱人工智能时代，其民事主体地位未被现行法律体系认可，侵权责任仍需追溯至人本身[5]。

在进一步转向追究行为人的责任时，便面临着第二个难题。生成式人工智能侵权往往涉及算法开发者、网络服务提供者、平台运营商、实际用户等多个行为主体，但算法“黑箱效应”使得算法开发者难以完全掌控模型生成行为，导致主观过错认定困难。网络服务提供者和平台运营商虽不直接参与生成，但其平台运营管理和提供服务行为可能助推侵权作品的传播，其注意义务的边界与过错认定标准也有待明确。而对实际用户来说，生成式人工智能的运行机制很难去完全掌握，同时其给出的指令也并非完全明确具体，既难以判断其主观侵权意图，也无法区分侵权结果究竟源于用户指令，还是算法的“黑箱效应”。可见，生成式人工智能侵权责任主体的认定需综合考量各方面因素，使得责任主体的认定难度进一步加大。

2.2. 因果关系难以证明

因果关系是联系侵权行为和损害结果的桥梁，也是认定侵权责任的关键环节，但在生成式人工智能侵权的认定过程中，往往会因其“算法黑箱”的技术特性而使得因果关系难以明确。这一特点决定了实际用户甚至算法开发者都难以明确生成式人工智能的整个思考和决策过程，以前文提到的著作权侵权为例，著作权侵权认定要求满足“接触 + 实质性相似”的双重标准，但实际上生成式人工智能的训练数据来源开放且海量，如果其输出的成品与特定作品构成相似，首先要证明该模型在训练中实际与特定作品存在实际接触，即便能从海量的训练数据中锁定大致来源，“算法黑箱”的特点也会导致无法精确追溯模型是否在数据处理中实际提取并使用了该作品的具体信息，而“实质性相似”的标准也会因算法对数据的碎片化借鉴而难以认定[6]。

2.3. 归责原则适用存在争议

归责原则是责任分配的核心依据与标准，传统的归责原则包括过错责任原则、无过错责任原则以及过错推定原则，但在生成式人工智能侵权的场景下，这三种归责原则都呈现出了一定的短板。首先就过错责任原则而言，被侵权人首先要证明侵权人存在过错，但受限于生成式人工智能“算法黑箱”的技术特征，实际上被侵权人往往难以提出有效证据证明侵权人存在过错，无法有效履行举证责任。即便是适用过错推定原则，人工智能提供者也通过证明自己已经履行了合理的注意义务而免除相应责任，长此以往反而会反向激励人工智能提供者在开发、设计、部署、制造过程中刻意减少相应成本，进而降低人工智能的透明度，加剧“算法黑箱”现象的发生[7]。而民法典规定无过错责任原则仅适用于高度危险作业等特定场景，如果盲目适用无疑会加重生成式人工智能服务提供者的责任，难以在科技创新和权利保障之间达到良性平衡。

3. 生成式人工智能侵权的法律规制的改善建议

3.1. 完善多层级的归责原则

针对归责原则难以统一适用的困境，可以依据生成式人工智能运行的不同阶段来分类规制。在数据

输入与模型训练阶段，算法开发者和提供者对系统的掌控力更强，能够预见数据侵权风险并采取防范措施，更适宜采用过错责任原则，通过客观化的注意义务标准认定过错；在内容输出阶段，生成式人工智能是依据用户提供的资料和指令进行输出，算法黑箱使过错认定失去客观基础，此时可引入过错推定规则，由掌握技术优势的服务提供者就自身已尽合理注意义务承担举证责任，以此缓解被侵权人举证不能的困境[8]。

还可参照“风险控制力强弱”与“损害严重程度”两项内容为标准构建风险归责框架，若服务提供者控制力强且侵权造成严重损害，则以严格责任为原则，免责条款仅限于受害人故意或不可抗力；若服务提供者控制力强而损害轻微，则推定其风险控制不当但允许反驳；若服务提供者控制力弱而损害严重，则应由被侵权人完成初步举证，随后转移举证责任[9]。通过建立层次化的责任配置原则来避免归责分歧，在权益保障与产业创新之间实现精细平衡。

3.2. 完善因果关系认定的技术化路径

针对“算法黑箱”的技术特性导致的因果关系证明障碍，可以从技术与制度两个方面来进行完善。首先，在技术层面上可以参考杭州互联网法院在“AI生成物著作权案”⁵中要求被告提供算法运行日志的做法设立算法运行档案强制留存制度[10]，要求开发者在设计之初即建立覆盖数据采集、模型训练到内容生成全流程的操作记录，系统留存数据集构成、参数调整与关键决策节点信息，当侵权纠纷发生时司法机关可依据档案还原过程，降低事实因果关系的证明难度。其次，优化“接触”要件的认定规则，若测试用户可通过特定提示词重现涉案作品或实质性相似内容，即可推定模型在训练中“接触”过该作品，举证责任转移至开发者。再次，从技术层面和专家审查两个方面构建起实质性相似的判断体系，在技术层面对特征相似度进行量化计算分析，在机制层面则由行业专家结合创作背景与商业价值进行定性判断，二者的结果相结合来提升侵权判定的科学性[11]。

3.3. 构建多元协同治理机制

面对责任主体多元化和责任关系复杂化的现实困境，可以着力构建多元协同治理机制。首先，明确生成式人工智能服务提供新型网络主体，在训练、交互与输出三阶段履行差异化注意义务，在训练阶段要审核数据来源合法性，在与用户交互阶段拒绝对恶意诱导指令的响应，输出期履行内容标识与投诉处理义务[12]。其次，推动开发者、运营者与使用者之间形成责任分担协议，通过事前约定明确各方在数据合规、算法安全与用户引导等方面的权责边界。再次，充分发挥公证制度作为预防性法律制度的独特价值，依托区块链技术搭建知识产权公证服务平台，实现从输入到产出的创作全流程证据保存制度，为数字证据难保存的问题提供新思路，在侵权发生后提供证据保全公证，为司法裁判提供关键证据支撑。通过多元主体的协同治理，可形成与链条式侵权特征相匹配的多元共治格局，在实现人工智能产业高质量发展与权利高效保障之间达成动态平衡。

4. 结语

当前，生成式人工智能侵权呈现出主体多元化、技术高度化与损害快速化等特征，传统侵权责任制度在主体界定、因果关系认定与归责原则适用等方面存在适用困难，难以在权益保障与技术创新之间实现有效平衡。

破解治理难题，既不能一味苛责服务提供者而抑制技术创新，也不能放任侵权乱象损害合法权益。本文构建了以风险控制力与损害程度为标尺的层次化归责体系，辅以算法档案留存、因果关系技术化认

⁵参见杭州市中级人民法院(2024)浙01民终10332号判决。

定及多元协同治理机制,通过设立多层次的归责原则,为生成式人工智能侵权的理论扩展提供支撑。

法律是社会关系的调节器,只有厘清各方主体权利义务边界,兼顾法律的稳定性与灵活性,方能为生成式人工智能产业塑造出合规有序的发展环境,实现技术进步与权益保护的动态平衡。

参考文献

- [1] 李锦华. 风险归责的体系化构建——人工智能侵权责任归责机制的再思考[J]. 数字法治, 2026(2): 129-145.
- [2] 程啸. 生成式人工智能侵害肖像权的侵权责任[J]. 法学论坛, 2026, 41(3): 5-16.
- [3] 白云霞. 人工智能生成物著作权侵权责任认定研究[J]. 知识经济, 2026(18): 154-159.
- [4] 杨钰涵. 生成式人工智能的“深度伪造”侵权责任重构——技术中立挑战与法律规制路径[J]. 产业创新研究, 2026(10): 40-42.
- [5] 谢黎伟, 潘圣. 生成式人工智能数据训练的著作权合规性困境和治理路径[J]. 山东科技大学学报(社会科学版), 2026, 28(3): 60-71.
- [6] 刘友华. 生成式人工智能版权侵权责任分层配置研究[J]. 法学杂志, 2026, 47(3): 67-87.
- [7] 张文睿, 刘颖泽. 生成式人工智能服务提供者归责原则的分层建构[J]. 河北法学, 2026, 44(7): 178-194.
- [8] 何泽昊. 从信息成本到信息生产: 人工智能侵权过错责任的理论证成[J]. 法商研究, 2026, 43(3): 65-79.
- [9] 章文丽. 人工智能应用的类型化侵权责任研究[J]. 上海师范大学学报(哲学社会科学版), 2026, 55(2): 58-69.
- [10] 上海市静安区人民法院课题组. 生成式人工智能服务提供者的侵权责任问题研究[J]. 数字法治, 2026(3): 181-193.
- [11] 黄术. 生成式人工智能辅助作品著作权侵权判定的理念塑造与规则调适[J]. 出版与印刷, 2026(3): 107-118.
- [12] 侯泽奕, 胡海涛. 生成式人工智能侵权责任的链条式规制: 从服务提供者中心主义到多主体责任分担[J]. 石家庄学院学报, 2026, 28(4): 93-102.