

生成式 AI 应用场景下版权侵权风险的 体系化规制

——以 DeepSeek 为视角

郭玉新

(清华大学 法学院,北京 100084)

摘要:DeepSeek 开启生成式 AI 的应用普及时代,丰富多元的应用场景增加了版权侵权风险规制难度。我国宜采取“前端宽松、过程控制、后端兜底”的进路,充分发挥著作权法利益平衡的功能,构建体系化的侵权风险规制方案。事前规范层面,基于公共利益以及产业政策目标设立生成式 AI 版权合理使用规则以满足人工智能发展需求,并设置梯度式豁免,厘清作品使用行为边界。事中约束层面,在“元规制”模式下,构建生成式 AI 版权市场自治体系,强化侵权风险识别与违法行为的行政规制,维护版权市场秩序,促进作品数据的高效流通,为版权人获益创造空间。事后救济层面,根据生成式 AI 版权侵权主体类型及侵权行为等因素对侵权责任的认定及承担进行类型化分析,明确法律责任分配,为版权人提供侵权救济通道。

关键词:生成式 AI;版权侵权;作品使用;体系化规制;DeepSeek

中图分类号:D923.41

文献标志码:A

文章编号:1672-7835(2025)03-0085-10

DeepSeek 以其高性能、低成本及开源模式等优势打破了人工智能技术壁垒,展现出强大的适配性和广泛的应用场景,在各领域迅速掀起“接入”浪潮。而随着 DeepSeek 的技术普及,在丰富的应用场景下,商业模式不断创新,版权侵权风险亦进一步“涌现”并呈现诸多新的特点。生成式 AI 开发者与服务提供者以及用户等各类主体面临多重侵权风险,生成式 AI 产业界与版权人之间的利益冲突有加剧之趋势。在此背景下,笔者对 2022 年以来美国、欧盟以及我国等国家或地区的生成式 AI 版权侵权案件进行不完全统计,共梳理案件 40 件,尝试厘清生成式 AI 版权侵权风险特征与演化趋势。据此分析和研判生成式 AI 版权侵权风险法律规制的主要挑战,结合我国人工智能产业发展需求与利益分配诉求,明确法律规制目标,突破单一的事后规制路径,构建事前规范、事中约束、事后救济相衔接的版权侵权风险应对方案,以期实现人工智能时代多元主体利益的精细化调整。

一、生成式 AI 应用场景下版权侵权规制难题与应然路向

(一) DeepSeek 浪潮下生成式 AI 应用与版权侵权风险

DeepSeek 推动生成式 AI 快速迭代,对于高质量作品的需求大大增加,版权侵权隐忧随海量作品的使用而凸显。通过采用强化学习算法、混合专家架构等全新技术路线,DeepSeek 突破了规模法则(scaling law)的限制,能够基于更少算力资源和成本并运用高质量数据训练高水平模型^①。同时,为实现提升模型性能、增强模型处理特定任务表现、匹配应用场景需求以及优化算法架构等目标,生成式 AI

收稿日期:2025-02-17

基金项目:国家社会科学基金重大项目(23&ZD159)

作者简介:郭玉新(1991—),男,内蒙古赤峰人,博士,助理研究员,主要从事知识产权法、数字法研究。

①李国杰:《DeepSeek 引发的 AI 发展路径思考》,《科技导报》2025 年第 3 期。

仍需持续利用数据进行优化训练。在 DeepSeek 技术溢出与示范效应下,生成式 AI 与作品权利保护之间的张力加大。从案例统计看,生成式 AI 版权侵权行为涵盖文字、美术、音乐、视听等作品客体类型,且随着模型技术升级以及文字作品数据减少,针对结构更为复杂的视听作品的侵权案件开始出现(2件)。

尽管 DeepSeek 拥有强大性能,但其仍存在“模型幻觉”问题,可能生成与特定作品构成“实质性相似”的内容。DeepSeek 创新性地采用“深度思考”与“联网搜索”模式,用户可选择其中一种或两种方式生成内容。上游使用者还可以通过“搭桥”“提示”等方式将大模型“换马甲”后为下游用户提供服务^①。而囿于模型训练或联网抓取的数据质量等因素限制,不同模式均可能生成侵权或危害性内容。从案例统计看,生成内容侵权的案件逐渐增多(达14件),体现出版权侵权风险随生成式 AI 的广泛应用而有所加剧。

DeepSeek 开源模式降低了人工智能的技术门槛,拓展了应用边界,使得版权侵权行为类型与侵权场景更加复杂。DeepSeek 可通过调用 API 接口、模型微调、蒸馏、嵌入等方式实现本地化部署与场景化应用,大量开发者或服务提供者结合市场需求开发专业模型(小模型、学生模型)并利用特定领域或行业的数据进行精细训练,进而为用户提供针对性服务。此外,随着生成式 AI 的个人化普及,用户可自建数据库并进行模型训练与生成内容,海量用户使用作品大大增加了侵权风险识别与监管的难度。在统计的案件中,权利人起诉侵犯复制权的案件数量最多(29件),其次是删除版权管理信息(12件)、制作衍生作品(9件)和侵犯发行权(8件)。在我国杭州互联网法院审理的“奥特曼”版权侵权案中,被告作为 AI 平台调用了第三方开源模型(Stable Diffusion)并进行技术整合,用户可以创建特定模型(LoRA 奥特曼模型)并通过特定模型生成内容(输入提示词后生成奥特曼相关图片)^②。

(二) 生成式 AI 版权侵权风险规制的主要难题

第一,生成式 AI 版权侵权问题随技术创新而产生,“以技术应对技术”具有优先性,但技术规制同样面临失灵问题。在数据输入阶段,海量数据的类型、格式、结构等特征千差万别,难以一一识别是否为受版权保护的作品,无法准确确认作品来源及权利状态。在数据训练阶段,数据的清洗、标注、预处理以及保存管理等技术环节难免产生偏差甚至错误,如删除版权管理信息、错误标注、数据泄露等。在内容输出阶段,尽管算法过滤技术可在一定程度上防止生成涉嫌侵权内容,但形式化的对比方式与以“实质性相似”为核心的版权侵权判定标准存在显著差异,难以保证内容过滤机制的效率以及准确性。因此,单纯依靠技术规制版权侵权风险不仅受到现有技术条件制约,同时还面临中小型企业技术能力以及成本、动力等因素限制。

第二,DeepSeek 的广泛应用衍生复杂的版权侵权风险,版权制度供给不足问题愈发突出。针对生成式 AI 使用作品的版权侵权挑战,现有研究从不同角度提出了诸多方案。有学者归纳出“非表达性使用论”“合理使用论”“法定许可论”“临时复制论”“版权侵权论”“总体国家安全观论”等五种观点并予以比较分析,提出赞成合理使用占相对多数^③。然而,既有规则面临不同程度的适用障碍问题。例如,我国现行《著作权法》中的合理使用规则难以直接适用于生成式 AI 情景,即便进行法律修订,也存在立法周期漫长等现实梗阻。概言之,持续的路线之争与制度适用障碍为人工智能产业发展带来了较大的不确定性。

第三,DeepSeek 的广泛应用对版权侵权纠纷司法裁判带来巨大挑战。一是侵权责任归责难。生成式 AI 版权侵权主体的法律性质尚未得到明确(如网络服务提供者、新型网络服务提供者、特殊责任主体、风险制造者),主观过错认定存在争议(过错原则抑或无过错原则),侵权行为较为随机且无法预测具体方式(生成内容因人因条件而异),因果关系判断受到模型类型、算法、应用方式及场景等复杂因素影响^④。二是侵权责任承担难。侵权人造成版权侵权应承担停止侵害、消除影响、赔偿损失等责任。然而生成式 AI 的“单向投喂”“持续迭代”等特征使得删除作品数据等救济请求难以落实,侵权行为可能

①苏宇:《大型语言模型的法律风险与治理路径》,《法律科学(西北政法大学学报)》2024年第1期。

②杭州互联网法院[2024]浙0192民初1587号民事判决书。

③易继明:《大模型语料训练合理使用问题研究》,《中国版权》2024年第6期。

④刘少军,聂琳峰:《人工智能生成内容的著作权法之辩》,《南昌大学学报(人文社会科学版)》2024年第1期。

无法消除;版权侵权损害赔偿计算亦是典型难题之一,尤其是如何评估作品在模型训练与内容生成中的贡献。

(三)全球生成式 AI 治理趋势下我国版权侵权规制路向选择

DeepSeek 引领生成式 AI 技术新一轮突破,推动全球人工智能治理体系与实践加速演化。人工智能全球治理呈现以分类分级为内核、软法与硬法动态衔接、监管主体跨域协同的特征^①。但与此同时,围绕人工智能技术与治理规则主导权的争夺日益激烈。美国、欧盟等国家或地区均表现出放松监管以加快技术创新的态势转向。例如,美国特朗普政府废除前任政府颁布的数十项人工智能监管政策以“扫除人工智能创新障碍”,包括取消 AI 开发者的安全测试强制报告义务等;同时计划投资 5 000 亿美元启动“星际之门”AI 基础设施建设,并渲染国家安全风险,加大对我国的技术封锁。欧盟通过《人工智能法》尝试继续扩大规则制定的“布鲁塞尔效应”,但严苛的监管措施与合规成本迟滞了其技术创新速度。DeepSeek 的技术突围增加了欧盟的 AI 雄心,新一届欧盟委员会有意改变监管政策,以缩小其与美国、中国的技术差距。此外,日本、韩国、印度、中国香港等国家或地区已开始或计划制定 AI 数据利用的豁免规则。可见,人工智能的全球治理进入“战国时代”。

DeepSeek 为我国完善人工智能发展战略与治理体系提供了新的契机。相较于美国“技术民族主义”的激进倾向与欧盟“规则理想主义”的保守思维,我国应保持定力,立足长远,以整体的视角来平衡技术创新及其风险治理,充分发挥著作权法利益平衡之功能。

在规制理念上,基于利益平衡与分配正义理念,构建风险共担、利益共享的激励相容法律规制体系,兼顾人工智能技术创新、产业发展与版权保护。一方面,高质量作品数据是支撑生成式 AI 技术创新的重要条件,而我国人工智能发展面临语料资源短缺、数据采集违法风险较高等阻碍。降低交易成本与侵权法律风险,构建高质量的语料供给机制是产业界的核心诉求之一。另一方面,版权人的核心关切是生成式 AI 未经授权使用作品且未作补偿构成侵权,生成内容可能对作品产生市场替代,挤压版权人的获益空间。若忽视版权人利益,过度限制其权利会产生“寒蝉效应”,不利于激发人类创作活力,对人工智能长远发展产生消极影响。因此,因应人工智能不同发展阶段,版权规则需要细致研判和回应各方利益诉求。

在实现路径上,围绕生成式 AI 版权侵权纠纷,既有以司法规制为核心的事后规制路径难以充分实现利益平衡。司法裁判的介入一方面回应了各方诉求,起到定分止争的作用,但亦可能引发新的争议。例如,部分法院尝试依据《生成式人工智能服务管理暂行办法》来认定生成式 AI 服务提供者的注意义务。但这一做法遭到质疑:监管规范难以成为著作权法上注意义务的设立依据,可能混淆公法与私法的边界。因此,事前事中事后相统一的体系化规制方案成为应然选择。具言之,事前阶段应以“前端宽松”为导向,解决作品专有权行使及其限制问题。根据社会公共利益以及产业政策目标,设立合理使用规则以满足人工智能发展需求。事中阶段以“过程控制”为导向,在“元规制”模式下,构建生成式 AI 市场自治体系,强化侵权风险的识别与违法行为的行政规制以保护版权人权益。事后阶段以“后端兜底”为导向,针对不同类型的侵权纠纷,通过司法救济机制明确生成式 AI 版权侵权的法律责任配置与利益分配。

二、事前规范:生成式 AI 使用作品行为的梯度豁免

严格的版权保护规则可能增加数据获取成本与合规成本,对生成式 AI 发展形成制约。纾解生成式 AI 应用场景下的版权侵权风险,首先需要从规范层面回应生成式 AI 使用作品行为的定性问题,基于合理使用规则构建梯度式的豁免框架。这有利于进一步激发我国生成式 AI 的发展活力。

(一)确立生成式 AI 版权合理使用规则

运用“卡—梅”框架分析,法律赋予版权人禁止他人使用作品的权利实质上是一种财产规则^②,他人

^①张欣:《人工智能治理的全球变革与中国路径》,《华东政法大学学报》2025 年第 1 期。

^②杨峰,刘先良:《卡—梅框架下我国排污权担保的规则配置研究》,《现代法学》2019 年第 5 期。

须得到版权人的授权许可并支付对价,否则版权人有权禁止使用。但财产规则与人工智能海量数据需求和高昂交易成本的现实情况存在天然冲突。而根据责任规则,可在必要情况下迫使法益“非自愿”转移,权利人则通过权威第三方作出的法定价格获取一定补偿。责任规则的本质是法益拥有者的定价权被法律强制买断^①。责任规则与生成式 AI 开发应用的现实需求、数据要素资源价值释放的政策目标以及知识产权开放共享之理念有着内在契合。著作权法中的合理使用便是一种责任规则。

确立生成式 AI 版权合理使用规则具有正当性、必要性与可行性。其一,DeepSeek 衍生丰富的 AI 应用场景与商业模式可能增加版权侵权风险,而合理使用作为侵权豁免事由,可消解生成式 AI 使用作品的侵权之虞,降低合规成本与创新成本,推动我国人工智能持续创新和广泛应用。其二,伴随 DeepSeek 引发的模型开源浪潮,包括版权人在内的广泛主体获得 AI 赋能,大大增加了社会整体福祉,这无疑强化了合理使用规则的正当性基础。其三,从国际趋势看,美国、欧盟、日本等国家或地区均通过变革合理使用规则来为其人工智能发展“松绑”。面对 DeepSeek 的冲击,OpenAI、Google 等巨头积极呼吁获得模型训练合理使用作品的豁免以保持其竞争优势。其四,我国确立人工智能合理使用规则已在学理与实践层面达成一定共识。例如,我国杭州互联网法院在“奥特曼”版权侵权案中提出,“在无证据证明生成式人工智能是为使用权利作品的独创性表达为目的、已影响到权利作品正常使用或者不合理地损害相关著作权人的合法权益等情形下,可以被认为是合理使用”^②。2025年3月13日公布的《人脸识别技术应用安全管理规定》第二条规定:“在中华人民共和国境内为从事人脸识别技术研发、算法训练活动应用人脸识别技术处理人脸信息的,不适用本办法的规定。”

不少学者提出,在《中华人民共和国著作权法》中新增“文本数据挖掘”例外条款,或在《中华人民共和国著作权法实施条例》中设置“为了 AI 学习和创作的使用”“为了进行数据挖掘,复制、存储他人作品以及将数据挖掘成果向公众提供”等规则。笔者认为,合理使用规则可表述为“为了人工智能开发和学习使用他人作品,但作者声明不许使用的除外”。为适应人工智能快速且未知的变化,合理使用规则的构建宜采用开放模式。其一,鉴于 AI 开发创新与产业应用的深度融合,合理使用不必明确排除商业目的,可结合具体场景与主体类型对豁免范围及程度作出差异化认定。其二,由于模型训练处于快速发展阶段,联邦学习等新型训练方式不断出现,合理使用无须明确排除特定作品使用行为。其三,为尊重和保护版权人合法权益,应在一定时期内为版权人设置“选择一退出”(opt-out)机制,即版权人可作出权利保留声明,未经其同意不得使用作品。

(二) 开源基础模型服务提供者的最大限度豁免

以开放协作共享为理念的开源运动颠覆了“大教堂式”自上而下的层级式创新模式,极大地拓展了创新疆域^③。而人工智能基础模型开源已超越传统开源软件(open source code)范畴而向开放资源(open resource)迈进。例如,DeepSeek-R1 模型采用 MIT 许可协议,开放了代码与模型权重,任何人可自由使用、修改、分发和商业化该模型,只需在所有副本中保留原始的版权声明和 MIT 许可。

开源基础模型适用合理使用规则建立在模型类型化分析的基础上。欧盟《人工智能法》第二条规定,除例外情形外,“本条例规定的义务不适用于根据免费且开源许可发布的人工智能系统”,即采取“技术开源+商业免费”的双重标准,“技术开源+有偿提供服务”等情形则不符合豁免要求。我国不少学者亦建议对免费且以开源方式提供人工智能的个人和组织给予减轻或免承担法律责任^④。笔者认为,在人工智能技术颠覆性发展的趋势下,商业模式处于不断变化和探索中,应以技术开源标准为核心,基于相对广泛的开源内涵,通过豁免技术开源的 AI 模型以激发创新活力。对于开源基础模型的认定,可根据开源许可证进行分析。开源许可证是 AI 开源社区运行和决定后续市场商业模式样态的主要

①凌斌:《法律救济的规则选择:财产规则、责任规则与卡梅框架的法律经济学重构》,《中国法学》2012年第6期。

②杭州互联网法院[2024]浙0192民初1587号民事判决书。

③张平:《开放创新的知识产权应用机制》,《知识产权》2024年第6期。

④杨建军,张凌寒,周辉,等:《人工智能法:必要性与可行性》,《北京航空航天大学学报(社会科学版)》2024年第3期。

“游戏规则”,体现出知识产权利用的不同样态^①。斯坦福人类中心人工智能研究院(HAI)根据开放程度将 AI 模型划分为七类。其中,属于开源范畴的包括三种:完全开源,即代码(包括推理代码、训练代码)、模型权重和训练数据均对外开放且无使用限制,如 GPT-NeoX 模型;上述要素开放但设有使用限制,如 BLOOM 模型;仅模型权重对外开放,如 Stable Diffusion 模型、Llama 2 模型^②。

结合合理使用“三步检验法”等标准,开源基础模型服务提供者的豁免范围可从“目的正当性”“市场影响”“作品影响”这三个方面予以分析,考量权重依次递减。其一,开源基础模型使用作品以技术创新为核心目的,能够带来公共利益增值并提高创新与生产效率,体现了合理使用所追求的公平和效率的法律价值目标。因此,当 AI 模型具有开源特征时,可推定其作品使用行为构成合理使用,是否为商业目的在所不问,除非有充足证据证明其具有主观恶意或显著的消极效果。其二,开源模型训练中可能涉及的作品使用行为均可纳入合理使用范畴,无须明确排除特定行为,如复制、转码、提取、汇编、解析、标注等,即不必预设某一行为绝对不构成合理使用。其三,基础模型使用作品并非意图生成与特定作品相似或相同之内容,而是形成新的生产工具。若生成内容与特定作品相似,但属技术问题导致或使用者诱导生成等非正常原因,模型服务提供者仍有获得合理使用豁免之可能,但豁免范围仅能延展至特定的生成内容场景,不能无限扩大。其四,当开源模型训练轻微影响作品的正常使用,如版权人作出权利保留声明后未被识别而实际使用了作品,或产生“过度拟合”“反复训练”等现象,模型服务提供者可基于主观无过错、已设立识别与优化机制等事由主张合理使用,并全面审慎评估其损害后果。

(三) 中小型生成式 AI 服务提供者的适当豁免

DeepSeek 催生海量应用层面的中小型生成式 AI 服务提供者,针对这类主体,合理使用的司法适用是宽松抑或严格,应根据具体情形进行具体分析。考虑到生成式 AI 应用场景下版权侵权风险相对增加,可将“市场影响”作为合理使用的核心考量因素,并结合“目的正当性”“作品影响”等因素综合研判行为的积极与消极效果。

一方面,生成式 AI 应用过程中,特定行业领域数据的定向“投喂”使得模型训练及生成内容对作品形成市场替代的风险相对较高,可能侵蚀版权人市场利益。例如,在 Thomson Reuters 诉 Ross Intelligence 案中^③,原告是具有全球影响力的大型法律信息提供商,其主张被告未经授权使用受版权保护材料进行训练和开发具有竞争关系的产品构成侵权。而法院在考量 AI 训练对作品市场价值的实质性损害程度后并未支持被告提出的合理使用抗辩。另一方面,尽管 DeepSeek 以其优异性能对各行各业带来变革,但作为基础大模型,其在具体应用场景中面临精确性、稳定性、技术与成本门槛等诸多问题。中小型生成式 AI 服务提供者则能够结合行业领域对基础模型作针对性调整、使用或开发专业模型,提供符合用户需求的场景化服务,推动技术落地并反馈技术创新环节。而一定规模和高质量的作品数据是保证模型获得优异性能的必要条件。

总体而言,为鼓励生成式 AI 应用普及,中小型生成式 AI 服务提供者对于作品的使用行为可构成合理使用,但应附加一定限制。一是鉴于中小服务提供者在义务履行能力、资源获取等方面的实际情况,可基于不同阈值,为中小服务提供者设置不同程度的豁免。对于严重影响版权人利益的行为予以否定性评价,如严格规制专门针对原作品风格的模型训练以及生成内容侵权行为。若生成内容侵权,则使用作品进行数据训练行为亦不构成合理使用,应支付费用。二是可为中小型生成式 AI 服务提供者的作品使用行为设置豁免期限(如 3—5 年),在此期限之内,相关主体可自由使用以鼓励创新,但应承担相对较高的注意义务。超出期限后可视经营规模、盈利能力等情况要求其付出合理成本以补偿版权人的利益损失。

^①辜凌云:《以许可证为核心的开源社区治理逻辑》,《知识产权》2024 年第 6 期。

^②Bommasani R., et al. “Considerations for Governing Open Foundation Models”, *Science*, 2024(386):151-153.

^③Thomson Reuters Enter. Ctr. GmbH v. Ross Intel. Inc., 1:20-CV-613-SB (D. Del. Feb. 11, 2025).

三、事中约束:生成式 AI 应用场景下版权侵权风险的过程规制

合理使用规则有利于加速人工智能技术创新与产业应用,但版权人的利益同样需要得到充分保障。生成式 AI 版权侵权危机的根源在于大规模的作品数据使用存在畸高成本而引发市场失灵^①。与事前阶段相衔接,事中阶段的规制重心在于从过程视角构建和维护版权市场交易规则及秩序,预防版权侵权风险并约束违法行为,进而促进作品数据的高效流通,为版权人分享人工智能发展红利创造空间。

(一)生成式 AI 版权侵权风险事中规制的目标导向

根据“卡—梅”框架,当交易成本过高时通常适用责任规则,交易成本较低时则适用财产规则,但两者并非绝对排斥,而是根据实际需求灵活选择。承前所述,生成式 AI 服务提供者可选择合理使用作品进行模型训练(当然亦可自愿支付费用或提供补偿,下文详述),除非版权人明确拒绝。这一规则设计在回应人工智能发展需求的同时为版权人保留选择余地,版权人可通过许可的方式行使权利。产业实践亦表明,囿于诉讼成本等因素,生成式 AI 服务提供者已与版权人进行合作,版权许可不断增加。例如,OpenAI 已与金融时报、美联社、Axel Springer 等内容提供商达成版权许可协议。

为生成式 AI 设置合理使用豁免,一个潜在的前提是生成式 AI 对于作品的使用行为应属版权专有权控制范畴,未经授权在理论上构成版权侵权。易言之,财产规则仍是生成式 AI 版权规则体系的重要基础。其一,具有独创性的人类创作作品是生成式 AI 持续创新的关键支撑。例如,DeepSeek 利用高质量数据进行模型训练,使模型预测的概率分布尽可能逼近实际数据的真实分布,模型的稳定性、鲁棒性和泛化能力得到提升,进而确保输出内容的准确性,避免产生内容偏见或歧视。版权人对其智力成果当然享有获取经济利益的正当权利^②。其二,著作权法的核心功能是激励功能,在制度上确认版权人可基于 AI 使用作品而获得收益,能够激发人们的创作热情。若从法理基础层面剥离版权人的请求权可能产生作品使用失控风险,进一步加剧利益失衡。

“卡—梅”框架在应用过程中衍生出了“管制规则”和“无为规则”两种规则类型。前者强调法律明确法益归属且允许私人转让,但代表国家的第三方权威会严格限定法益转让的法定条件;后者则在法定情形下否认或取消了一项利益获得法律救济的法益资格,因而无涉权益归属、交易和定价问题^③。管制规则通过政府监管将市场交易纳入规则范围,从而遏制侵权风险。在“合理使用+版权许可”的利益分配模式下,即便生成式 AI 服务提供者获得合理使用豁免或得到版权人许可,其仍应满足法定要求,否则会引致监管介入。诸如超出合理使用范畴的作品使用行为、生成内容与特定作品构成“实质性相似”以及因商业模式创新而引发新的侵权问题,均需要通过监管予以应对。

总之,事中阶段法律规制的目标并非完全消弭风险,而是通过构建预防和疏导机制,维护版权市场机制的有效运行,实现产业发展与权利保护的动态平衡。

(二)生成式 AI 版权侵权风险事中规制的路径构建

面对生成式 AI 应用引发的版权市场失灵,应充分发挥有效市场与有为政府的共同作用,推动多元主体协同共治。对此,可构建“元规制”(meta-regulation)的整体性规制框架。所谓元规制是以初级规制为对象的次级规制,即以社会自我规制为对象的政府规制。其中,市场主体自我规制为首要和主导方式,政府的再规制为必要补充^④。DeepSeek 引领生成式 AI 几乎为各个行业带来巨大变革,自发市场秩序处于生成或快速变化阶段,法律规制面临高度的信息不对称。元规制模式可激发市场自治的灵活性优势,降低规制成本,契合包容审慎的监管理念;政府的再规制则能进一步克服自我规制存在的私利干

①孙山,张雯雯:《生成式人工智能预训练中权利限制制度的选择与建构》,《科技与出版》2024年第7期。

②Senftleben M. “Generative AI and Author Remuneration”, *International Review of Intellectual Property and Competition Law*, 2023(54): 1535-1560.

③凌斌:《法律救济的规则选择:财产规则、责任规则与卡梅框架的法律经济学重构》,《中国法学》2012年第6期。

④黄文艺,孙喆玥:《论互联网平台治理的元规制进路》,《法学评论》2024年第4期。

扰、权利滥用以及责任逃避等问题^①。

1. 生成式 AI 版权市场自我规制

生成式 AI 版权市场自我规制的目标是优化市场机制在分配版权利益、促进作品流通利用以及预防风险等方面的功能,其关键在于从数据获取及输入、数据训练、内容输出及使用等方面构建灵活的版权市场合规机制。

在作品获取方面,“合理使用+版权许可”的利益平衡模式要求生成式 AI 服务提供者与版权人等主体积极履行义务以破解作品获取难题。其一,生成式 AI 服务提供者可通过合理使用的方式获取作品,但应承担作品来源信息的识别与保留义务,以确保作品来源合法。合规重点是在技术与管理等方面建立识别与筛查机制以避免侵权,并且不得删除具有权利来源识别功能的版权管理信息(CMI)。其二,市场自我规制旨在降低交易成本、推动形成高效的作品许可机制。实践中,中小型生成式 AI 服务提供者作为数据需求方可能面临被大型内容平台拒绝许可、缔约成本过高等困境^②。作为供给方,掌握海量作品资源的大型内容平台在版权许可市场中具有较强的议价能力,应基于竞争性义务“公平、合理、无歧视”地进行作品许可,不得滥用知识产权。针对分散的海量版权人,可强化著作权集体管理组织的功能,或设立市场化的新型管理组织,代表版权人开展版权许可与利益分配。

在作品使用方面,DeepSeek 在垂直行业领域广泛的私有化部署及开放接口使得数据泄露风险随之增加。一方面,模型训练完成后,采集、存储的海量作品数据可能因保管不当而泄露,对此,模型开发者应及时采取删除、保存或移交等措施;另一方面,在模型优化及应用场景下,企业或用户会自行“投喂”数据,从而直接或间接产生数据泄露。例如,通过提示词注入攻击(prompt injection attack)等方式诱导模型生成原始数据,以及通过多次获取数据片段进行“数据拼图”来非法获取原始数据。对此,生成式 AI 服务提供者应承担风险提示义务,提示用户上传数据并用于模型训练的风险。

在生成内容方面,DeepSeek 的广泛应用使得生成内容的不可控性增加,版权侵权风险随之加剧。对此,生成式 AI 服务提供者可能需要承担内容过滤、侵权投诉处理、算法优化、消除重复作品数据等义务。

2. 自我规制的政府再规制

“元规制”理念旨在避免行政部门过早、过度干预人工智能版权利益格局的市场形塑过程,在协同共治理念下,推动生成式 AI 版权治理主体从政府向企业、行业组织的横向拓展以及规制权力的纵深配置。

第一,对生成式 AI 版权市场自我规制予以监督和约束。不同类型的生成式 AI 服务提供者承担差异化的合规义务,面临不同的合规成本,其中数字平台等大型市场主体具备“准立法权”“准司法权”“准行政权”而成为市场自治规则的主导者。然而,平台公共性与私利性的内在冲突始终存在。有鉴于此,政府规制的重点是防止数字平台在市场自治中的权力异化风险,同时提升整体版权合规水平。具体可通过发布版权合规指引、合规监督检查等方式,促使大型平台构建版权合规机制并引导中小型主体根据自身情况履行合规义务,提升市场自治的透明性、公平性。同时,建立合规激励机制,例如,可参考《企业知识产权管理规范》国家标准的经验,为完成版权合规的生成式 AI 服务提供者提供资金补贴等。

第二,充分利用 DeepSeek 等开源生态所具有的信息发现及传递机制,提升多元主体互动与协同效率。生成式 AI 版权侵权风险的主要挑战之一便是法律规制面临较为显著的信息不对称。因此,政府部门需要强化不同治理主体之间的协同联动,实现信息互联互通与资源整合。行政机关可与版权人、研发者以及行业组织等主体合作,围绕作品的使用和内容生成,建立具有高拓展性的风险监测机制,动态追踪并及时干预侵权风险。

第三,生成式 AI 版权侵权风险具有场景多元、样态复杂、风险随机等特征,需要法律规制的及时快速响应与精细化处置。对此,积极运用 DeepSeek 赋能智慧监管,开发部署版权侵权风险预警模型,智能识别

^①王海洋:《生成式 AI 训练数据的法律风险及其元规制》,《浙江社会科学》2024 年第 9 期。

^②孙靖洲:《人工智能训练的版权困境及其出路:模块化许可机制探析》,《知识产权》2024 年第 11 期。

高危风险领域和监管对象,提升风险预警与处置的科学性、准确性。行政机关还可通过提示、约谈、指导等软性规制工具,要求人工智能企业按照相关规定自行整改,降低行业“野蛮生长”带来的负面影响。

四、事后救济:生成式 AI 应用场景下版权侵权责任的类型化分析

(一)生成式 AI 版权侵权责任分配规则

DeepSeek 能够广泛进入垂直领域的企业级实际应用,衍生各种类型的商业模式,为多元化用户提供内容生成(文本、音视频)、内容设计(视觉、文案)与内容分发(搜索、推荐)等服务。在此背景下,对生成式 AI 服务提供者进行分类分级规制具有现实必要性。

就主体类型而言,存在风险视角与功能视角的区分标准。欧盟《人工智能法》根据不同风险级别将人工智能系统划分为不同类别并设定相应义务,但存在规制目标泛化、规制方式滞后、规制标准严苛等缺陷,难以适应生成式 AI 多形态、多场景的现实要求^①。功能视角则强调生成式 AI 应用环节差异:一是“简易分层”,即区分模型开发与模型提供,前者可获责任豁免;二是“弱分层”,即对人工智能(产品、服务)研发者与提供者的责任作有限区分;三是“强分层”,即在“弱分层”的基础上进一步区分不同行业或场景,形成“基础模型—专业模型—服务应用”之界分^②。

就主体性质而言,生成式 AI 服务提供者是否构成著作权法意义上的网络服务提供者是学理和司法实践中的争议焦点。生成式 AI 服务提供者根据用户指令生成内容,这种“算法创作”模式与传统网络服务存在显著差别。特别是 DeepSeek 将推理过程显示出来,生成内容并非原始数据的简单堆砌和呈现,而是类似人脑思考的“创作”。由此,生成式 AI 服务提供者的定性存在“新型网络服务提供者”^③“特殊责任主体”^④等不同观点。

综上所述,生成式 AI 版权侵权责任的认定首先应区分生成式 AI 基础模型与应用服务两个层面,并视情况对应用服务层作进一步细分。进而,再结合生成式 AI 产品或服务特征、行业领域与具体场景等因素对版权侵权责任主体性质及责任分配进行个案分析。

第一,以 DeepSeek 为代表的基础模型服务提供者需就其直接向最终用户提供生成内容服务时的侵权行为承担过错责任。根据侵权法的一般原理,在法律没有特别规定的情况下,生成式 AI 服务提供者应当承担过错责任。在司法实践中,法院可根据生成式 AI 服务提供者是否履行特定注意义务来判断其主观过错,反之,生成式 AI 服务提供者可通过“避风港”规则主张责任豁免。基础模型服务提供者具有较为显著的技术创新目的及效果,加之开源模式有利于技术扩散、产业应用并促进公共利益,不宜采取严格责任,以防止对我国人工智能创新造成阻碍。故而,基础模型服务提供者的注意义务在范围与程度方面应进行限缩。

第二,当基础模型为其他主体所应用并向最终用户提供生成内容服务时,除非基础模型存在与版权侵权具有重大关联的严重技术风险或其他严重情形,否则基础模型服务提供者不对下游环节的侵权行为承担责任。原因在于,基础模型对于终端生成内容的风险控制力随着应用链条的延伸而不断衰减,难以对风险进行预判和审查;要求基础模型提供者对下游主体承担风险减轻义务或安全保障义务,可能大大加重其负担,抑制基础模型创新活力。

第三,应用层的生成式 AI 服务提供者应作为版权侵权责任的主要承担主体,但需根据主体类型和侵权行为等因素作差异化规制。一方面,此类主体应负有相对较高的版权注意义务,例如作品获取与使用时的权利识别、作品数据的安全保障以及生成内容的过滤、提示、投诉处理、删除以及主动筛查等;另一方面,法院可通过公平原则与诚信原则,综合考量各方举证能力等因素,要求应用层的生成式 AI 服务提供者承担相对较高的举证责任。作出上述安排的原因在于:通过接入 DeepSeek 等开源基础模型,应

①丁晓东:《人工智能风险的法律规制——以欧盟〈人工智能法〉为例》,《法律科学(西北政法学报)》2024年第5期。

②苏宇:《大型语言模型的法律风险与治理路径》,《法律科学(西北政法学报)》2024年第1期。

③冯晓青,沈韵:《生成式人工智能服务提供者著作权侵权责任认定》,《法治研究》2025年第1期。

④姚志伟:《生成式人工智能服务提供者在私法上的法律性质》,《上海交通大学学报(哲学社会科学版)》2024年第12期。

用层的生成式AI服务提供者直接面向最终用户,对于版权侵权风险的形成、扩散具有重要影响,同时其已具备独立的风险控制能力与获益条件,应当承担与之相匹配的法律责任。当然,一刀切地加重应用层生成式AI服务提供者的责任可能对人工智能产业发展带来负面影响。对此,应进一步结合行业领域、具体场景、保护法益等作出细化调整。例如根据经营规模、技术能力、用户规模等因素区分大型与中小型生成式AI服务提供者,明确责任差异;考量新闻、搜索、科研、医疗、社交、生物识别、自动驾驶等不同行业领域的特征,合理评估版权侵权损害。

(二)生成式AI版权侵权的损害赔偿与利益补偿

在“卡—梅”框架下,生成式AI版权侵权行为可能同时包含对财产规则和责任规则的违反,但两种情形下的责任承担并不相同^①。

第一,财产规则下,生成式AI服务提供者生成内容的版权侵权损害赔偿应依次根据版权人的实际损失、侵权人的所得利益、许可使用费、法定赔偿四种计算方式予以确定。在笔者统计的案件中,囿于实际损失与侵权获利均难以计算,原告通常提出实际损失、被告获利以及法定赔偿等多项主张且数额较高。法院在适用法定赔偿时应充分平衡各方利益。我国广州互联网法院在“奥特曼”版权侵权案中综合考量了涉案作品的知名度、被告采取的补救措施及其效果、侵权的影响范围等因素^②。杭州互联网法院则区分生成式AI不同应用场景,分类分层界定侵权责任,并充分考虑生成式AI新兴商业模式发展需求^③。沿此路径,可建立法定赔偿分档规则,根据权利类型及特点(如作品类型与使用数量、作品市场价值)、侵权主体类型(如开源基础模型、应用模型)、主观过错(过失、故意)、侵权行为严重性(如生成内容数量、传播范围)以及补救措施等因素确定法定赔偿区间及具体数额。

第二,对于生成式AI生成内容的版权侵权,应审慎适用惩罚性赔偿责任。生成式AI产业界与版权人以及社会公共利益之间的利益平衡格局尚未发展成熟,轻率适用惩罚性赔偿可能激化各方矛盾。而生成内容的技术特征(如系统内呈现)亦使得侵权损害程度相对有限。因此,仅当生成式AI对权利人造成严重损害时方可考虑惩罚性赔偿。例如,以营利为目的故意模仿特定权利人作品风格大量生成内容并传播。总之,财产规则下的惩罚性赔偿旨在使侵权人的利益状态恢复到甚至低于侵权行为实施前的水平,实现预防和遏制侵权的主要目标,而非单纯追求惩罚效果,其适用条件应受到严格限制。

第三,根据责任规则,生成式AI服务提供者可通过构建利益补偿机制来获得责任减免。一方面,生成式AI服务提供者合理使用作品,但其仍可选择以一定方式为版权人提供合理补偿,这一举措有助于激励版权开放而扩大作品数据供给;另一方面,当作品使用超出合理使用范畴,如使用已作出权利保留声明的作品,或生成内容构成侵权时,生成式AI服务提供者可以“事后补偿”的方式弥补版权人经济利益,以此“换取”不承担停止侵害等侵权责任之豁免并允许继续使用作品^④。具体而言,一是通过直接的货币补偿、提供差异化或定制化的AI产品与服务(如在生成内容次数、时长等方面予以优待)等方式补偿版权人。二是与版权人达成收益分配协议。对此,可借鉴目前内容平台“版权管理系统”(Content-ID)模式,让版权人获得二次收益(如流量分成、广告分成)。三是合理平衡权利人之间的利益分配。对于合作作品,合作作者通过协商一致来行使著作权和分配收益,协商不一致的可根据贡献程度公平合理分配收益。对于职务作品,一般职务作品著作权由作者享有,其通常可以允许生成式AI使用作品,除非对单位利益造成显著影响,收益分配则遵从双方之约定。对于委托作品,若未作约定或未订立合同,著作权属于受托人,其可自行决定版权收益之分配。

第四,探索建立政府主导的利益补偿机制,为遭受侵权损害的版权人提供补偿。第一,可根据生成式AI服务提供者的经营规模、商业模式、技术能力等收取一定数额的资金(如以版税方式);第二,政府积极运用产业政策、财政政策、金融政策等扶持生成式AI发展,并划拨专项资金;第三,通过社会组织、

①魏远山:《“卡梅框架”视角下我国知识产权惩罚性赔偿制度的完善》,《西部法学评论》2023年第2期。

②广州互联网法院[2024]粤0192民初113号民事判决书。

③杭州互联网法院[2024]浙0192民初1587号民事判决书。

④刘云开:《人工智能生成内容的著作权侵权风险与侵权责任分配》,《西安交通大学学报(社会科学版)》2024年第6期。

行业协会、金融机构等筹集资金^①。当然,上述补偿机制是人工智能发展初期为促进技术创新与应用而采取的权宜方案,其需要辅以适当的监管机制以防止市场主体逃避责任,产生劣币驱逐良币的不利后果,确保真正具有技术创新或商业模式创新能力的企业得到支持。

结语

DeepSeek 引领我国生成式 AI 走向一个新的关键节点,在这场深刻的变革中,技术风险与治理挑战相伴而生。正如当代科学家斯蒂芬·沃尔夫勒姆(Stephen Wolfram)所提示的,人工智能时代真正的威胁是人类在符号洪流中丧失定义规则的能力。DeepSeek 热潮下,构建人工智能风险治理体系是把握技术发展航向的关键。其中,生成式 AI 与版权制度的冲突及协调问题成为核心议题,需要审慎处理。在技术创新及应用的颠覆性变化阶段,产业界与权利人之间的利益平衡格局尚处形成过程,有必要构建体系化的法律规制方案以实现人工智能发展与作品保护以及社会公共利益的动态平衡,并为踟蹰与摇摆中的人工智能版权规则与司法实践提供路径参考。

Toward a Systematized Paradigm for Regulating Copyright Infringement Risks in Generative AI Applications: A Case Study of DeepSeek

GUO Yuxin

(School of Law, Tsinghua University, Beijing 100084, China)

Abstract: With the advent of DeepSeek, generative artificial intelligence has entered a phase of unprecedented proliferation, engendering an array of application scenarios whose richness and variability exacerbate the intricacies of copyright infringement governance. To navigate these complexities, it is incumbent upon China to adopt a stratified regulatory schema characterized by leniency at the inception stage, vigilant procedural oversight, and remedial safeguards ex post facto—thereby actualizing the equilibrium-seeking mandate of copyright jurisprudence. At the preliminary regulatory tier, provisions for fair use tailored to the generative AI context ought to be instituted, rooted in overarching considerations of public interest and national industrial policy imperatives. These provisions should operationalize a gradient-based exemption structure responsive to the heterogeneity of involved actors, while concurrently demarcating the legitimate boundaries of content utilization. At the midstream stage of regulation, under the auspices of a meta-regulatory framework, a self-governing copyright ecosystem should be architected within the generative AI sector, one that augments the capacity for infringement risk detection and fortifies administrative oversight mechanisms. Such a system would serve to preserve normative order in the copyright marketplace and facilitate the fluid dissemination of copyright-encumbered data, thereby enhancing revenue-generating opportunities for rights holders. In the aftermath of potential infringements, liability allocation should be subjected to a typological analysis predicated upon the categorical nature of the infringing party and the specific modalities of the infringing conduct. This approach would clarify the apportionment of legal responsibility and furnish efficacious avenues of redress for aggrieved copy right owners.

Key words: generative artificial intelligence; copyright infringement; utilization of protected works; systemic regulatory framework; DeepSeek

(责任校对 曾小明)

^①黄玉烨,杨依楠:《论生成式人工智能版权侵权“双阶”避风港规则的构建》,《知识产权》2024年第11期。