

生成式人工智能服务提供者的版权注意义务

——以“通知 - 删除”规则失灵为视角

李佳琪

青岛大学法学院, 山东 青岛

收稿日期: 2026年8月9日; 录用日期: 2026年9月2日; 发布日期: 2026年9月9日

摘 要

传统“通知 - 删除”规则以侵权内容的可定位性和可移除性为操作前提。生成式人工智能的技术特性使这一前提发生动摇: 在输入端, 训练数据被转化为模型参数后无法通过删除操作予以移除; 在输出端, 生成内容的概率性与非确定性使服务提供者难以对侵权内容进行精准定位和阻断。这两个层面的障碍共同导致以“特定侵权内容定位 - 删除”为核心的传统操作逻辑发生功能性失灵。但“失灵”不等于通知机制整体失效——通知仍然能够提供重要的版权风险信息。真正需要转化的是通知的法律效果: 从“删除触发机制”转化为“风险控制触发机制”。据此, 应当在明确服务提供者新型网络服务提供者定位的基础上, 以风险控制能力为核心标准重构版权注意义务体系, 在训练数据、生成运行、通知处置、动态测试四个阶段分别设置差异化的注意义务, 并在过错认定中综合考量风险可预见性、技术可行性与防控成本等因素, 实现版权保护与技术创新之间的动态平衡。

关键词

生成式人工智能, 服务提供者, 通知 - 删除规则, 注意义务, 版权侵权

Copyright Duty of Care of Generative AI Service Providers

—From the Perspective of the Failure of the Notice-and-Takedown Rule

Jiaqi Li

School of Law, Qingdao University, Qingdao Shandong

Received: August 9, 2026; accepted: September 2, 2026; published: September 9, 2026

文章引用: 李佳琪. 生成式人工智能服务提供者的版权注意义务[J]. 服务科学和管理, 2026, 15(5): 1007-1015.
DOI: 10.12677/sssem.2026.155117

Abstract

The traditional “notice-and-takedown” rule operates on the premise that infringing content is locatable and removable. The technological characteristics of generative AI undermine this premise: at the input stage, training data are transformed into model parameters and cannot be removed through deletion operations; at the output stage, the probabilistic and non-deterministic nature of generated content makes it difficult for service providers to precisely locate and block infringing content. These two levels of obstacles jointly lead to the functional failure of the traditional operational logic centered on “locating and deleting specific infringing content”. However, this “failure” does not mean that the notification mechanism as a whole has become ineffective—notifications can still provide important copyright risk information. What truly needs to be transformed is the legal effect of notifications: from a “deletion-triggering mechanism” to a “risk-control-triggering mechanism”. Accordingly, on the basis of clarifying the service provider’s status as a new type of internet service provider, the copyright duty of care system should be reconstructed with risk control capability as the core standard, with differentiated duties of care established across four stages—training data, generation operation, notification response, and dynamic testing—and with factors such as risk foreseeability, technological feasibility, and prevention costs comprehensively considered in fault determination, so as to achieve a dynamic balance between copyright protection and technological innovation.

Keywords

Generative Artificial Intelligence, Service Provider, Notice-and-Takedown Rule, Duty of Care, Copyright Infringement

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. 问题的提出

生成式人工智能是指具有文本、图片、音频、视频等内容生成能力的模型及相关技术，旨在通过学习大量数据的分布，捕捉数据的内在特征和模式，在保持数据分布一致性的基础上进行新内容的生成[1]。

“通知－删除”规则是网络版权治理的基石性制度，凭借技术中立、事后救济、责任豁免的制度逻辑，成为全球网络平台版权治理的通用范式。我国《信息网络传播权保护条例》¹确立了本土化的“通知－删除”框架，《民法典》²第 1195 条进一步从民事基本法层面固化其效力，使之成为网络服务提供者免责的核心依据。传统规则的有效运行依托于两个基本前提：其一，网络平台仅承担中立、被动的技术存储与传输功能，不参与内容创作与加工；其二，网络侵权内容以独立文件形态存储于服务器，具备可定位、可删除、可重复防控的特征。

生成式人工智能的规模化应用，从根本上动摇了这两个前提。与传统网络平台单纯传播已有作品不同，生成式 AI 依托海量数据完成深度学习，能够基于用户指令自主生成全新内容。有学者指出，生成式人工智能行为在技术上的特殊性，对侵权责任带来了一系列影响，包括行为主体的复杂化、加害行为的

¹ 《信息网络传播权保护条例》2006 年 5 月 18 日中华人民共和国国务院令第 468 号公布，根据 2013 年 1 月 30 日《国务院关于修改〈信息网络传播权保护条例〉的决定》修订。

² 《中华人民共和国民法典》2020 年 5 月 28 日第十三届全国人民代表大会第三次会议通过，自 2021 年 1 月 1 日起施行。

智能化、损害后果的不确定、因果关系的多元化和过错认定的新颖化[2]。这一技术变革冲击了“通知-删除”规则的适用根基，引发制度功能性失灵。从输入端来看，训练数据经算法拆解、特征提取与参数迭代后内化为模型参数，原始侵权作品无法单独定位与删除；从输出端来看，AI生成内容具备概率性、随机性特征，侵权内容无固定形态与生成路径，无法通过传统拦截、删除方式实现长效防控；从主体层面来看，生成式AI服务提供者兼具技术服务提供者与内容间接塑造者的双重属性，打破了传统网络主体二元划分结构，有学者将其界定为介于传统内容提供者与服务提供者之间的“新型网络服务提供者”[3]。

需要明确的是，传统规则失灵并非意味着通知机制彻底失效。服务提供者因部署具有固有风险的人工智能系统而持续开启风险，且具备事实上的风险控制能力而应尽相应注意义务。权利人发出的侵权通知，依然能够精准指向具体、特定的版权侵权风险，打破平台“不知情、无过错”的免责状态，成为平台过错认定的重要前提。当前司法实践已逐步摆脱僵化的“唯删除论”，不再以平台是否删除侵权内容作为唯一归责标准，而是转向审查平台是否尽到合理版权注意义务，形成了从结果导向的侵权处置到行为导向的风险防控的裁判转型。

基于此，本文以“通知-删除”规则失灵为切入点，剖析传统制度的适用困境，以风险控制能力为核心重构生成式人工智能服务提供者的版权注意义务体系，完善司法过错认定与利益平衡规则，为人工智能时代网络版权治理提供适配的制度路径与裁判参考。

2. 传统“通知-删除”规则在生成式AI场景下的适用困境

传统“通知-删除”规则的适用困境，本质是固有制度逻辑与生成式人工智能技术特性的根本性冲突。这种冲突体现在输入端、输出端、主体端三个维度，原有“定位-删除-防控”的运行链条难以完整运行，亟需完成制度功能的迭代重构。

2.1. 输入端：训练数据无法定位与删除

传统“删除”操作的核心前提是侵权内容以独立、具象、可检索的文件形式存储。生成式AI的训练过程彻底改变了数据的形态。平台用于模型训练的各类版权作品，不会以原始文件形式留存，而是经算法反向传播、特征压缩、迭代计算后，转化为海量分布式参数弥散于整个模型架构之中。训练数据的核心价值是让模型习得创作规律与语义特征，而非复刻原始作品。这种参数化存储模式导致单一侵权作品无法精准定位，更无法单独清除其侵权影响。

目前学界与产业界探索的机器遗忘、模型编辑、增量训练等技术，普遍存在精准度不足、算力成本高昂、整改效果难以验证等问题，尚不具备规模化适用的产业可行性。唯一能够彻底消除侵权影响的全量重训模式，因成本过高而难以落地。因此，传统“通知-删除”规则所预设的“定位-删除”操作链条，在生成式AI的输入端面临根本性障碍。

2.2. 输出端：生成内容无法确定和持续控制

传统“通知-删除”规则的另一个预设是：侵权内容具有确定性，同一侵权内容可以被反复识别和定位。权利人发出通知后，服务提供者不仅可以删除该内容，还可以通过技术手段防止其再次出现。生成式AI的输出机制使这一确定性不复存在。

AI的输出是基于模型习得规律实时计算生成的全新内容，而非对既有作品的简单复制。同一提示词在不同模型版本、不同交互场景下会产生差异化输出，侵权内容无固定形态、无固定生成轨迹。用户可通过微调提示词、更换创作场景、拆分创作指令等方式，规避平台的常规防控措施，生成与原版权作品实质性相似但表达形式不同的侵权内容。

有学者指出，生成式人工智能服务提供者在技术上无法完成“通知-删除”规则下的断开链接或删除等必要措施。这使得传统规则的“即时制止”与“重复防控”两大核心功能面临重大挑战。广州互联网法院在“奥特曼案”中指出，被告虽然在诉讼过程中采取了一定的技术措施，但未能证明这些措施可以有效防止被侵权作品的再次生成。

2.3. 主体端：平台定性不清导致规则适用僵局

“通知-删除”规则的适用以服务提供者被定性为“网络技术服务提供者”为前提。生成式 AI 服务提供者的双重属性使这一认定陷入困境。形式上，平台仅提供算法工具与智能生成服务，不直接产出固定内容，符合技术服务提供者的定位；实质上，平台通过筛选训练数据、设计算法逻辑、搭建内容过滤机制、迭代模型参数，深度决定生成内容的创作风格与合规水平，对侵权内容的产生具有实质性影响。

主体定位的模糊性引发规则适用的僵局。如果平台被定性为内容提供者，则其可能承担直接侵权责任，无法通过“通知-删除”免责；如果被定性为技术服务提供者，则可以适用该规则，但技术上的障碍又使“删除”无法真正实现。有研究指出，由于相关规范未能对“提供内容”与“提供服务”加以明确区分，致使平台在主体属性和归责事由上出现法源选择和适用的双重难题。同时，《生成式人工智能服务管理暂行办法》³赋予平台的公法内容监管义务，与私法侵权责任认定规则无法直接衔接，公法合规与私法追责的错位进一步加剧了裁判的困难。

2.4. 从“删除”到“处置”的功能转换

综合上述三个层面的分析，可以得出一个基本判断：传统“通知-删除”规则以特定侵权内容的可定位、可删除和可持续控制为基础，在生成式 AI 场景下，其核心操作逻辑发生了功能性失灵。

需要区分三个层次来理解这种“失灵”。第一，通知并未失灵——权利人的通知仍然能够传递重要的版权风险信息。第二，处置并未失灵——平台在接到通知后仍然可以采取屏蔽提示词、限制生成功能、调整模型、封禁用户账号等多种措施。第三，真正失灵的是“通知后直接删除特定侵权内容”的传统操作逻辑——“内容定位→删除→防止再次传播”这一链条在输入端因数据不可定位而断裂，在输出端因内容非确定性而难以闭环。

因此，需要将通知的法律效果从“删除触发机制”转化为“风险控制触发机制”。通知的功能不再要求平台从模型中“删除”某个已不存在的独立文件，而是要求平台在获知特定版权风险后，在其技术能力范围内采取合理的风控措施。这一功能转化，为后文讨论注意义务的转向奠定了制度基础。

3. 版权归责的机制转向：以注意义务替代传统删除规则

3.1. 通知机制的功能转化

在传统避风港规则中，通知的功能是触发删除义务。在生成式 AI 场景下，虽然“删除”难以实现，但通知本身仍具有不可替代的法律价值。

具体而言，有效通知可以发挥三重功能。其一，风险告知功能——通知使平台从“不知”状态进入“应知”状态，明确了特定版权风险的存在。其二，义务升级功能——在接到通知后，平台的注意义务程度相应提高，不能再以“对侵权行为不知情”为由主张免责。其三，措施指引功能——通知中载明的具体信息(如权利作品名称、相关提示词等)可以为平台采取针对性措施提供方向。

服务提供者因部署具有固有风险的人工智能系统而持续开启风险，且具备事实上的风险控制能力而

³《生成式人工智能服务管理暂行办法》(国家互联网信息办公室、国家发展和改革委员会、教育部、科学技术部、工业和信息化部、公安部、国家广播电视总局令第15号, 2023年)。

应尽相应注意义务。通知的上述三重功能，正是将抽象的注意义务转化为具体行为要求的制度接口。

3.2. 注意义务作为过错认定标准的正当性

在我国侵权法体系中，注意义务并非独立的侵权责任构成要件，而是过错认定的客观判断标准。《民法典》第 1165 条第 1 款确立的过错责任原则，要求被侵权人证明行为人具有过错。而过失的判断采取客观标准，即考察行为人是否达到了“合理的人”在相同情况下应当达到的注意程度。

在生成式 AI 版权侵权的场景下，注意义务的功能就是将“合理的人”这一抽象标准具体化为可操作的行为要求。有学者指出，服务提供者过错的本质是其违反了交往安全义务或注意义务，在认定标准上应采取动态的“合理人”标准[4]。过失认定仍应坚持以“理性人”标准为基石，但需将注意义务的适用范围从单纯的人类行为合理性拓展至“人机交互系统的整体安全性”，以契合技术特性与风险结构。

注意义务标准相对于“删除”标准具有明显的规范优势。其一，动态适配性——注意义务可以根据技术发展水平动态调整，不会因为“删除不可行”而陷入制度空白。其二，过程性定位——不同于删除规则僵化的结果导向要求，注意义务属于“方式性义务”，仅要求平台在自身技术能力范围内尽到合理防控职责，不苛求绝对杜绝所有侵权结果，精准契合 AI 概率性生成的技术本质。即便平台服务中出现个别侵权内容，只要其已全面履行合理注意义务，即可排除主观过错、免于承担侵权责任。

3.3. 法定义务与注意义务的衔接

《生成式人工智能服务管理暂行办法》等行政监管规范为平台设定了数据合规、风险排查等法定义务。公法义务是民事注意义务的重要参考，但二者不能等同：公法以维护公共秩序为目的，私法注意义务以保护私权为核心，规制客体与责任标准均有显著差异。平台违反行政管理义务，不当然构成民事侵权过错。但“不当然构成”过于笼统，未能回应一个核心问题：平台完全遵守公法规范时，在民事侵权诉讼中能否以及何种程度上作为已尽注意义务的抗辩？对此需从以下层次进行类型化分析。

（一）合规类型的差异化效力

其一，实体性合规——低度抗辩效力。《暂行办法》第 7 条要求使用具有合法来源的数据和基础模型，平台若履行了数据来源合规审查，该事实可作为已尽合理注意义务的初步证据。但公法审查以形式合规为主(如授权文件齐备性)，不要求对数据中隐含的侵权内容进行实质全面筛查，故完全合规不等于充分注意，法院仍须结合版权风险的明显程度综合判断。

其二，程序性合规——中高度抗辩效力。公法中涉及风险识别与内部治理的程序性义务(如建立投诉举报机制、开展自查)。平台若系统落实此类义务，表明其构建了风险防控的制度化框架，可在较大程度上证明其尽到了“合理人标准”的注意义务。尤其在侵权内容生成路径难以精准预判时，程序性合规的完备性能够有力反驳“应知”的过错推定。但若制度形同虚设，权利人仍可举证推翻抗辩。

其三，宣示性合规——无独立抗辩效力。纯属倡导性的原则性规定(如“鼓励创新发展”)，因无可操作的行为指引功能，合规与否对过错认定不产生实质影响。

（二）领域关联性与动态权衡

公法合规抗辩的效力还取决于合规行为与版权侵权之间是否具备“领域关联性”：若合规仅涉及算法备案等与版权保护无关的程序，不得以此主张尽到版权注意义务；反之若合规恰好覆盖了涉案侵权的风险防控维度，法院应予以充分考量。此外，同一侵权路径若在合规后多次出现且未有效阻断，表明合规措施存在执行漏洞，抗辩效力应显著削弱。

综上，公法合规在民事侵权诉讼中的证据效力，应遵循“领域关联 - 类型区分 - 动态权衡”的判断框架，既尊重公法监管价值，又避免“合规即免责”的简单化推理，确保注意义务认定回归以风险控制

能力为核心的综合判断标准。

3.4. 风险控制能力作为核心判断标准

风险控制能力是注意义务配置与过错认定的核心标尺，遵循“风险开启、权责一致”的基本侵权法理。生成式 AI 服务提供者是版权侵权风险的开启者与最优控制者，完整掌控模型架构、数据来源、算法逻辑、过滤机制与服务运营全流程，相较于权利人与用户，具备更强的风险识别、预防与处置能力，理应在自身控制能力范围内承担匹配的注意义务[5]。

有学者明确指出，注意义务的边界应当与服务提供者对风险的控制能力相匹配，控制能力越强，注意义务越重；控制能力越弱，注意义务越轻。服务提供者对风险的控制受到现有科技发展阶段的客观限制，因此控制能力与注意义务之间的关系首先受“现有技术条件”制约[6]。

在司法实践中，需综合以下核心要素构建弹性化的过错判断规则：(1) 风险可预见性，以行业普遍认知判定侵权风险是否可预判，知名 IP、热门作品的风险预见标准更高；(2) 技术可行性，以行业通用技术水平为基准，不苛求顶尖、超高成本的防控技术；(3) 平台控制能力，根据训练、生成、交互等不同环节的控制力差异化定责；(4) 损害程度与发生概率，批量、系统性侵权的追责标准更为严苛；(5) 平台获利情况，商业化有偿服务的注意义务标准显著高于免费服务。最终实现平台风险、控制能力与法律责任的有机统一。

4. 生成式 AI 服务提供者版权注意义务的体系化构造

基于 AI 全流程运行逻辑，结合平台不同阶段的风险控制力差异，可构建覆盖数据训练、生成运行、通知处置、动态测试的四阶段差异化版权注意义务体系。

4.1. 训练数据阶段的注意义务

训练数据阶段平台控制力最强，需承担最高标准的源头合规义务。

(1) 数据合法性审查义务

《生成式人工智能服务管理暂行办法》第 7 条要求，服务提供者应当使用具有合法来源的数据和基础模型，涉及知识产权的不得侵害他人依法享有的权利。平台需建立标准化的数据溯源、权属核查与合规筛选机制，优先采用合法授权或符合合理使用条件的数据集，主动剔除权属存疑的版权作品。该义务为方式性义务，仅要求平台遵守行业审慎合规流程，因技术局限导致少量侵权数据进入训练集的，不当然认定为过错。

(2) 数据记录与举证协助义务

平台需完整留存训练数据来源、合规审查记录、模型训练日志等核心资料，纠纷发生时在保护商业秘密的前提下为权利人举证提供必要协助。该义务的功能在于降低版权纠纷中的证明困难，服务提供者若能证明其已对数据来源进行了合理审查并保留相应记录，可作为已尽合理注意义务的重要证据。

4.2. 生成运行阶段的注意义务

生成运行阶段侵权风险高发且存在天然技术局限，平台需承担中等强度的常态化防控义务。

(1) 高风险版权对象的识别与过滤义务

针对知名度高、侵权风险明确的作品，平台需建立专项识别与拦截机制，当用户输入明确指向侵权复刻的提示词时及时阻断生成路径。该义务仅聚焦高风险场景，不苛责平台对海量普通作品进行全量筛查。广州互联网法院在“奥特曼案”中即基于该形象的极高知名度，认为服务提供者应当能够预见侵权风险。

(2) 合规提示与引导义务

平台应在服务界面公示 AI 内容合规使用规则，对疑似侵权的生成场景弹出风险提示。该义务虽不能直接阻止侵权，但可起到警示和引导作用。

4.3. 通知处置阶段的注意义务

平台接收有效通知后，常规义务升级为专项处置义务。

(1) 基于合理通知的风险处置义务

对于要素齐全、侵权路径可复现的有效通知，平台需采取多元化处置措施，包括侵权提示词屏蔽、相似变体指令拦截、违规用户管控、模型局部参数优化等。有学者主张，当特定提示语与侵权内容存在稳定复现关系时，可将修改相关模型参数纳入必要措施范围[7]。

(2) 比例原则与长效防控

处置措施应与侵权频次、损害轻重相适应，既杜绝消极不作为，也避免过度处置损害正常服务。必要措施的选择应考虑技术可行性，不能要求平台采取当前技术水平下无法实现的措施。平台需将已确认的侵权路径纳入风险数据库，持续优化过滤规则，降低重复侵权概率。

4.4. 动态风险测试阶段的注意义务

针对模型迭代和用户规避带来的新型风险，平台需承担主动预防的动态测试义务。对已涉诉或高风险的版权对象，定期开展多场景测试，覆盖提示词变体、创作场景切换等常见规避方式，核验现有过滤机制的有效性。针对测试发现的问题，迭代优化算法模型与风控词库。该义务以已知高风险场景为限，测试频次与整改力度应与风险等级匹配，避免过度负担。

5. 版权注意义务的司法适用

5.1. 有效通知的认定标准

有效通知应包含以下要素：权利人身份信息和联系方式；作品权属证明；能够指向特定侵权生成路径的描述信息；初步证明“特定提示词与侵权内容之间存在稳定复现关系”的证据；要求采取的具体措施。对于存在轻微形式瑕疵的通知，可允许权利人补充完善，不宜直接否定效力。批量维权通知可认可概括性清单举证。

5.2. 注意义务违反的判断因素

法院判断是否违反注意义务，需综合权衡以下因素：风险的可预见程度——以同业服务提供者参照；技术可行性与防控成本——以行业通用技术水平为基准，不苛求超越现有技术条件的措施；平台的实际控制能力——控制力越强、义务越高；损害的严重程度与发生概率；平台的获利情况——有偿服务的注意义务标准高于无偿服务。

5.3. 利益平衡与裁判导向

注意义务为方式性义务而非结果性义务，平台不承担杜绝一切侵权的结果责任。法院应通过动态化、类型化的裁判规则适配技术迭代，区分不同运营阶段、服务类型、平台规模差异化裁判，避免“一刀切”式归责。注意义务的制度目标既包括有效保护版权人合法权益，也包括为人工智能技术创新和产业良性发展预留充足空间。

6. 域外主要法域 AIGC 版权治理的实践与比较

生成式 AI 版权治理是全球性难题。美国、欧盟、英国、日本基于各自制度传统与产业政策，形成

了差异化的应对路径。梳理这些实践，可为本文以注意义务为核心的治理方案提供比较法上的正当性支撑。

6.1. 美国：司法主导下的合理使用调适

美国以判例法为主导推进规则形成。2025 年，Bartz v. Anthropic 与 Kadrey v. Meta 两案均认定 AI 训练构成合理使用，但前者对训练数据来源作出关键区分：合法购买扫描的数据可用于训练，而来自盗版数据库则不行。后者引入“市场竞争”理论——AI 生成内容即使不直接替代原作品，仅形成市场竞争关系也可能构成侵权，扩张了责任边界。此外，2026 年提出的 CLEAR 法案要求开发者在商业发布前向版权局披露训练数据中的版权作品，但尚未通过。整体上，美国以合理使用为 AI 训练提供安全港，但数据来源合法性和输出端侵权风险正成为新的争议焦点。

6.2. 欧盟：透明度与选择退出的立法先行

欧盟采取成文法先行的强监管路径。《数字化单一市场版权指令》⁴第 4 条赋予权利人“选择退出”(opt-out)文本与数据挖掘的权利；《人工智能法案》⁵第 53 条进一步要求通用 AI 模型提供者公开训练数据摘要(2025 年 8 月生效)，但摘要仅披露数据来源和性质，不涉及具体作品细节，权利人验证困难。欧洲议会 2026 年决议建议增设法定报酬权、建立集体管理机制，并设置最高全球营业额 4%的罚款。欧盟模式规则明确，但 opt-out 的可验证性和跨域适用效果有待检验。

6.3. 英国：政策摇摆中的特定选择

英国目前仅允许非商业性文本与数据挖掘，商业 AI 训练须取得许可。2024-2025 年公众咨询中，88% 回应者支持“所有情形均须取得许可”，政府倾向的“选择退出”方案仅获 3%支持，政策推进受阻。英国尚处于十字路口，其最终选择将影响欧洲治理格局。

6.4. 日本：“非享受目的”例外的独特路径

日本《著作权法》第 30 条之 4 规定⁶，“非以享受作品中表达的思想或情感为目的”的利用无须经许可。AI 训练若仅为学习创作规律，可纳入该例外。但该条款的解释边界存疑——当生成内容与原作高度相似时，训练目的是否仍属“非享受”存在争议。日本模式为 AI 产业提供了宽松预期，但权利人保护不足的问题日益凸显。

6.5. 比较与启示

上述实践可归纳为三种模式：美国“司法合理使用”模式(灵活但不确定)、欧盟“透明度+选择退出”模式(明确但验证难)、日本“非享受目的”模式(宽松但保护弱)。尽管路径各异，但各国均面临传统“通知-删除”规则在 AI 场景下失灵的共同困境，治理重心正从“事后删除”转向“事前预防与过程合规”。

对本文的启示有三：其一，欧盟的透明度要求与本文“数据记录与举证协助义务”功能一致；其二，美国对数据来源合法性的关注支撑了“数据合法性审查义务”的设定；其三，各国普遍转向“行为导向”的风险防控，印证了本文以注意义务替代删除规则的核心命题。中国在借鉴域外经验时，应立足本土产业实际，在版权保护与技术新闻寻求动态平衡——这正是本文四阶段差异化注意义务体系的制度目标。

⁴Directive (EU) 2019/790 of the European Parliament and of the Council of 17 April 2019 on copyright and related rights in the Digital Single Market and amending Directives 96/9/EC and 2001/29/EC.

⁵Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act).

⁶日本《著作权法》(昭和四十五年法律第四十八号，2018 年修订)第 30 条之 4。

7. 结语

传统“通知-删除”规则以特定侵权内容的可定位、可删除和可持续控制为基础，在生成式人工智能场景下，其核心操作逻辑发生了功能性失灵。本文的核心命题是：这种失灵不等于通知机制整体失效，也不意味着平台可以免除所有责任。真正需要调整的是通知的法律效果——从“删除触发机制”转化为“风险控制触发机制”。

这一转化的实质，是从“结果导向”的侵权处置逻辑向“行为导向”的风险防控逻辑的转换。法院不再追问平台有没有删除特定侵权内容，而是审查在已经知道或应当知道某种版权风险的情况下，平台是否采取了与其风险控制能力相适应的合理措施。这一转换既承认了技术局限的客观存在，又通过注意义务的设定保持了法律对服务提供者的行为约束。

以注意义务为中心的版权侵权责任框架，需要在训练数据、生成运行、通知处置、动态测试四个阶段进行体系化的制度构造。这一构造以服务提供者的风险控制能力为核心判断标准，以综合考量风险可预见性、技术可行性和防控成本为判断方法。通过这一体系化建构，可以在“通知-删除”规则失灵的情况下，为版权保护提供替代性的制度保障，同时为技术创新预留必要的发展空间。

参考文献

- [1] 陶建华, 赫然, 刘偲, 主编. 生成式人工智能[M]. 戴琼海, 主审. 北京: 清华大学出版社, 2025.
- [2] 吉聪聪. 生成式人工智能服务提供者侵权责任研究[D]: [硕士学位论文]. 济南: 山东财经大学, 2026.
- [3] 程啸. 生成式人工智能服务侵权不应适用网络侵权规则[J]. 比较法研究, 2026(1): 125-137.
- [4] 沈森宏. 论生成式人工智能服务提供者过错的认定[J]. 现代法学, 2024, 46(6): 133-147.
- [5] 姚佳. 生成式人工智能服务提供者的注意义务[J]. 法学论坛, 2026(3): 17-31.
- [6] 王若冰. 论生成式人工智能侵权中服务提供者过错的认定——以“现有技术水平”为标准[J]. 比较法研究, 2023(5): 20-33.
- [7] 王利明, 包丁裕睿. 论“通知”规则在生成式人工智能作品侵权中的类推适用[J]. 比较法研究, 2025(4): 1-13.