

AIGC时代新闻传播的伦理法律风险及 分层治理

魏宝萱

北京印刷学院新闻传播学院, 北京

收稿日期: 2026年6月22日; 录用日期: 2026年7月28日; 发布日期: 2026年8月4日

摘 要

生成式人工智能(AIGC)正从辅助工具演变为新闻生产网络中的准行动者, 其介入新闻场域后衍生多重治理难题, 如人机权责模糊消解传统把关机制, 算法生成逻辑弱化新闻人文底色, 训练数据侵权与AI生成物权属争议带来系统性版权纠纷, 深度伪造内容加剧网络信息失序等。本文以新闻专业主义与行动者网络理论为分析框架, 将AIGC对新闻业的根本性挑战归结为事实性根基动摇与公共性逻辑消解双重维度, 从而拆解伦理与法律风险的内在关联, 并从立法监管、媒体内控、从业者能力重塑、技术企业合规、受众素养提升五大主体出发构建分层协同治理体系, 为智媒时代新闻业的高质量发展提供学理参照与实践路径。

关键词

生成式人工智能, 新闻传播, 新闻专业主义, 伦理困境, 媒介治理

The Ethical and Legal Risks of News Transmission and Hierarchical Governance in the AIGC Era

Baoxuan Wei

School of Journalism and Communication, Beijing Institute of Graphic Communication, Beijing

Received: June 22, 2026; accepted: July 28, 2026; published: August 4, 2026

Abstract

Generative artificial intelligence (AIGC) is evolving from an auxiliary tool to a quasi-actor in news production networks. After its intervention in the news field, it has arisen multiple governance

challenges, such as blurring human and machine rights responsibilities to resolve traditional shut-down mechanisms, algorithmic generation logic to weaken the background of news literacy, training data infringement and AI-generated property rights disputes to lead to systemic copyright disputes, deep falsification of content to exacerbate network information disorder, and so on. This paper uses journalism professionalism and actor network theory as an analytical framework. It reduces AIGC's fundamental challenges to the news industry to the dual dimensions of factual fundamental instability and public logic resolution, thus dismantling the intrinsic association between ethical and legal risks. It builds a hierarchical and collaborative governance system starting from the five main subjects of legislative regulation, media internal control, practitioner capability reshaping, technology enterprise compliance, and audience literacy enhancement, providing a theoretical reference and practical path for the high-quality development of the news industry in the intelligent media era.

Keywords

Generative AI, News Communication, News Professionalism, Ethical Dilemmas, Media 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

伴随大语言模型与多模态生成技术的迭代突破,生成式人工智能已全面介入新闻采集、写作、可视化与分发各环节。新华社“快笔小新”、各大媒体的AI虚拟主播等实践,标志着新闻生产正从人主导的单主体模式向人机协同的双主体模式转型。技术中性视角下,AIGC为新闻生产带来效率跃升与形态革新,但其算法逻辑与数据规则的先天缺陷,同步催生虚假新闻、版权侵权、人文缺失等多重风险。梳理现有研究可见两个明显不足:其一,多数研究或聚焦伦理失范,或专注法律规制,缺少将二者纳入统一分析框架的系统性审视;其二,对策建议普遍偏向宏观原则,完善法律法规、提升职业素养却未能给出分主体、可落地的操作路径。本文尝试弥补上述缺憾,先厘清AIGC介入后新闻传播业态的核心特征,再从伦理与法律双维度拆解风险的内在机理,最终分层构建适配行业现实的多主体协同治理方案。

2. AIGC 时代新闻传播的业态新特征

2.1. 内容生产：从劳动密集型转向人机协作型

依托海量多模态数据的预训练,AIGC可自动完成数据整理、热点抓取、初稿撰写等重复性环节。AIGC技术冲击了新闻生产的生产关系,使新闻机构能够自动化生成新闻稿件[1]。面对财经快讯、体育简讯、天气播报等标准化新闻,AI能在数分钟内产出完整稿件,显著压缩生产周期。从媒介技术演进视角审视,这并非简单的效率提升,而是新闻生产劳动分工的结构性重构。机械性、程式化的采编环节被技术替代,从业者得以将认知资源集中于深度调查、价值判断与意义阐释等高阶劳动。这一转型的挑战在于,如何在效率增益与专业坚守之间建立动态平衡。

2.2. 受众分发：从大水漫灌到算法分众推送

AIGC联动用户画像与算法推荐,依据受众阅读时长、内容偏好、互动行为生成定制化推送,打破传统媒体的统一播发模式。分众化传播提升了信息触达效率,精准匹配不同圈层的信息需求。但需警惕的

是, AI 时代算法推荐技术是“信息茧房”产生的重要驱动因素[2], 算法以最大化用户停留时长为核心优化目标, 天然偏好高情绪唤醒度、强冲突性的内容, 这不仅为信息茧房埋下技术伏笔, 更形成从注意力指标驱动到猎奇内容优先再到公共议题边缘化的恶性循环, 公共信息供给不足则进一步削弱媒体公共价值, 降低自身议程设置能力。

2.3. 新闻呈现：从单一文本到多模态叙事融合

AIGC 打通文本、音频、图像、虚拟数字人的一体化生产链路, 可一键生成新闻短视频、交互 H5、VR 沉浸式报道等多种形态。多模态产品适配短视频平台的传播逻辑, 丰富新闻叙事形式, 拓宽主流媒体在新型终端的影响力。然而, 多模态生成技术在降低视听内容生产门槛的同时, 也同步降低了其深度伪造的制作难度。当伪造视频足以以假乱真, 新闻业赖以立足的影像即证据的认知基础将面临根本性动摇。

3. AIGC 时代新闻传播的伦理挑战

3.1. 生产主体重构：人机权责边界模糊与专业主义弱化

传统新闻场域中, 记者、编辑是内容生产的唯一主体, 完整遵循选题、采访、核实、编审的专业化流程。拉图尔的行动者网络理论打破了传统的主客体间的二元对立状态, 将人与非人存在者都视作网络中的行动者, 在这个网络中的“人”失去了本体论上的先在性[3]。AIGC 进入新闻编辑部后, 正经历一个双向驯化过程。从业者试图将 AI 界定为辅助工具, 以维持自身核心行动者地位, 但 AI 也反向转译着记者行为, 使其逐渐适应 AI 的输出格式、接受其效率逻辑, 甚至产生内容审美上的算法趋同。这一过程催生三重风险: 其一, 低门槛的关键词生成新闻模式弱化了实地走访、事实核查的职业惯习, 部分从业者直接照搬 AI 生成内容, 放弃一线调查; 其二, 长期依赖 AI 产出标准化内容, 可能使深度写作与思辨分析能力慢性退化; 其三, AI 批量生产的同质化稿件挤占深度特稿的版面资源与受众注意力, 可能导致基层采编岗位的结构性失业, 形成数字劳工困境。

3.2. 把关机制失效：后真相语境下的虚假信息泛滥

卢因的把关人理论在 AIGC 时代遭遇结构性冲击, 传统的多层级人工审核链条被技术便利性解构。这一冲击可从三个层面加以审视, 技术层面, 大模型存在幻觉缺陷, 会以高度逻辑自洽的文本形式捏造不存在的数据、信源与事件, 这类虚假内容因其仿真性而难以被常识判断识别; 主体层面, 普通网民与自媒体均可零成本生成看似专业的新闻内容, 把关主体从专业机构泛化至全员; 分发层面, 推荐算法以情绪唤醒度为核心指标, 而猎奇、冲突、恐慌类虚假内容恰好在这些指标上表现优异, 形成算法偏好假新闻的悖论。三重机制叠加, 使 AIGC 时代的虚假信息呈现出高仿真性、高传播性、低识别度的全新特征, 持续消解媒体机构的公信力积累。

3.3. 叙事逻辑异化：数据理性对人文关怀的置换

新闻客观性是新闻理论中极其重要的命题之一[4], 要求报道应在客观事实与人文温度之间寻求平衡。而 AIGC 依托统计规律生成内容, 其运作逻辑本质上是数据的关联与预测, 而非意义的理解与阐释。在灾难、事故、民生困境等报道中, AI 仅能机械罗列伤亡数字与财产损失数据, 缺乏对受灾者个体经历与情感诉求的共情性呈现。这种数据化叙事将鲜活的生命经验抽象为统计类别, 极易引发受众对报道的疏离与抗拒。更深层的风险在于模板化生产对新闻多样性的侵蚀, 当 AI 学习的是海量文本中的平均表达, 其生成内容必然趋近统计学意义上的均值, 新闻写作由此陷入公式化窠臼, 独特视角与创新叙事被系统性消解, 媒体与公众之间的情感联结也随之断裂, 弱化公众对媒体的情感认同。

4. AIGC 时代新闻传播的法律挑战

4.1. 训练数据与生成物的双重著作权困境

AIGC 的法律风险贯穿输入端与输出端两个环节,呈现结构性叠加特征。输入端,大模型训练需抓取海量网络文本、新闻作品、书籍等内容,未经著作权人许可的大规模复制与商用训练。输出端,现行法律明确规定著作权保护对象为自然人的独创性智力成果,纯 AI 生成内容无法获得版权保护。这意味着媒体使用 AI 稿件后,面临着尴尬的权属真空,既无法主张自身权利,又难以防范他人无偿使用。此外,模型训练数据常包含未经脱敏处理的个人信息,生成内容时可能无意识泄露,衍生隐私侵权风险。

4.2. 深度伪造与网络信息生态失序

信息失序领域权威学者克莱尔·沃德尔提出的信息失序三分框架有助于精准界定 AIGC 带来的不同风险。第一层为误讯(misinformation),指无意的错误信息,AIGC 幻觉生成的虚假内容多属此类;第二层为谬讯(disinformation),指蓄意制造的虚假信息,深度伪造领导人讲话、伪造当事人影像等属之;第三层为恶讯(malinformation),指基于真实信息但用于制造伤害,如将公开数据重新语境化以攻击特定对象。AIGC 同时加剧了三种形态,误讯因海量 AI 生成内容涌入而激增;谬讯因深度伪造门槛骤降而泛滥;恶讯因数据整合能力增强而更趋隐蔽。多重叠加造成网络空间的真伪难辨与信息冗余,严重扰乱公共信息秩序。现行平台审核技术在快速甄别 AI 伪造内容方面仍有较大滞后,事后溯源取证难度大,违法成本与危害后果严重不对称。

4.3. 技术依赖引发的从业者合规风险

人工智能主导的机器信源生成过程通常依据已有的数据模板和运算规则,事件原本的复杂细节可能会被简化,从而导致新闻内容脱离于事件的原始语境[5]。从业者过度依赖 AI 衍生双重风险,既属于职业伦理问题,也会触发明确法律责任,双重风险叠加加剧媒体合规危机。职业层面,记者省略实地采访、交叉信源核实等核心流程,直接发布 AI 生成内容,一旦内容失实,媒体机构需承担名誉侵权、不实报道的民事责任;行政监管层面,要求标注 AI 生成内容,大量媒体未建立内部规范、未显著标识 AI 创作部分,存在被网信部门行政处罚的风险。合规事故频发会持续降低媒体官方身份的权威性,损耗制度真实性。

5. AIGC 新闻传播的分层协同治理策略

5.1. 国家立法层面:完善法律法规,明确技术使用红线

针对当前法律空白,建议推进以下立法完善。在著作权法领域,参照域外经验探索文本与数据挖掘例外规则的适用条件,同时明确人机协同作品的版权归属标准,可参照人类创造性贡献的实质性程度进行分层界定。在深度伪造规制领域,出台专项管理条例,强制要求 AI 生成视听内容添加不可篡改的数字水印与元数据溯源标识。在行业监管层面,完善 AIGC 新闻服务准入规范,明确媒体使用生成式 AI 的内容审核义务,建立 AI 虚假新闻的梯度处罚机制。

5.2. 媒体机构层面:建立内控机制,规范人机协作流程

建立科技伦理准则、规范及问责机制,形成人工智能伦理指南,建立科技伦理审查和监管制度,明确人工智能相关主体的责任和权力边界,推动 AIGC 迈向“负责任”,已经成为 AIGC 技术伦理治理的一种必然趋势[6]。媒体应制定统一的 AIGC 使用伦理规约,核心条款应包括明确 AI 仅作为辅助工具,深度调查与重大主题报道必须以实地采访为核心;设立独立的 AI 内容审核岗位,所有 AI 生成稿件须经人

工二次事实核查后方可发布；建立强制标识制度，在版面及视频显著位置标注 AI 参与程度，如全文生成、辅助写作、数据可视化等。同时，建立常态化内部培训机制，围绕 AI 技术原理、新闻伦理与法律法规开展定期轮训，防止从业者专业能力的隐性退化。

5.3. 从业者层面：重塑核心能力，确立人机协同的职业认知

面对 AIGC 的冲击，新闻从业者的核心竞争力将从信息采集转向价值判断、意义阐释与伦理守护。具体而言，需构建三项新型核心能力，一是算法审计素养，能够识别并校正训练数据中的隐性偏见，主动辨别模型幻觉与虚假数据；二是合成内容核查技能，掌握反向图像搜索、语义一致性交叉分析等 AI 内容验证方法；三是人机边界判断力，能够根据报道题材与伦理敏感度，独立判断 AI 介入的适当程度。从业者应将精力聚焦于深度调查、独家叙事与人文阐释等 AI 不可替代的核心领域，以差异化竞争力应对同质化冲击。恪守新闻伦理和社会责任的专业素养，面向人工智能 AIGC 大模型不仅要做到“守土有责”，更要做到“开疆扩土”，开辟面向大模型训练的可信数据集和数据服务新阵地，提供决定大模型核心能力和价值观的内容供给与知识供给，抢占 AIGC 时代舆论引导、思想引领、文化传承、服务人民的传播高地[7]。

5.4. 技术企业层面：研发治理工具，践行技术向善

人工智能企业应承担与其技术能力相匹配的治理责任。在事前预防端，优化大模型训练的数据授权机制，建立合规授权的素材数据库，避免未经许可抓取版权内容；研发并嵌入 AI 内容数字水印与生成溯源技术，实现全流程可追溯。在事中监测端，开发大模型幻觉识别工具与深度伪造检测系统，并向新闻机构开放使用。在事后分发端，配合媒体优化推荐算法，弱化单一流量指标的考核权重，增加深度报道、公共议题内容的推荐占比，从技术架构层面缓解算法茧房效应。

5.5. 社会受众层面：普及媒介素养，构建社会共治网络

虚假信息治理的末端防线在于公众的辨别能力。媒体、教育机构与网信部门应联合开展全民媒介素养提升工程，面向不同年龄层普及 AIGC 技术原理与 AI 虚假内容识别方法，重点培养信源追溯与多方交叉验证的信息消费习惯。同时，畅通虚假信息便捷举报渠道，建立从举报到核实最终及时辟谣的快速响应闭环，鼓励受众从信息的被动接收者转变为内容生态的主动监督者，形成全社会共治共享的良性生态。

6. 结论

AIGC 作为颠覆性媒介技术，在提升新闻生产效率、丰富传播形态、创新分发模式的同时，也因其实算法逻辑与数据规则的先天缺陷，衍生出虚假信息泛滥、版权归属悬置、人文关怀缺失等多重治理难题。本文系统拆解了新闻伦理与法律风险的内在机理，并构建了涵盖立法、媒体、从业者、技术企业、受众五大主体的分层协同治理框架。研究的未尽之处在于，文中所提治理路径仍属框架性建构，其实际效度有待实证检验。后续研究可从面向全国地方融媒体开展问卷调查、开展新闻编辑部田野观察的方式来刻画基层从业者 AIGC 使用困境并深描真实人机协同工作流程与隐性冲突这两方面聚焦展开。此外，还可以采用实验法测试受众对 AI 新闻、人类原创新闻的信任感知差异，最终推动治理方案场景化、精细化落地。

参考文献

- [1] 赵子墨. AIGC 背景下新闻生产的变革与重构[J]. 采写编, 2024(7): 34-36.
- [2] 张省, 蔡永涛. 算法时代“信息茧房”生成机制研究[J]. 情报理论与实践, 2023, 46(4): 67-73.

- [3] 李捷. 行动者网络理论视阈下的机器人主体性问题研究[J]. 齐齐哈尔大学学报(哲学社会科学版), 2017(7): 34-36+43.
- [4] 张佳雯, 龚砚庆. 刍议新闻客观性的中西分野与中国路径[J]. 新闻爱好者, 2025(12): 77-79.
- [5] 杨奇光, 张宇. “机器信源”视角下新闻真实认识论危机及其检视[J]. 中国出版, 2025(18): 11-17.
- [6] 蔡津津. AIGC 时代新闻舆论工作新阵地——面向大模型的可信训练数据集与服务能力建设[J]. 中国传媒科技, 2023(10): 79-83.
- [7] 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.